

SIGNAL GENERATORS

Use and calibration

[Based on E. M. E. R. (Tels. Z. 025) Issue 1]

INTRODUCTION

1. The purpose of this Instruction is threefold:—

(a) To serve as an introduction to the use and operation of signal generator calibration equipment.

(b) To examine the validity of signal generator measurements as at present carried out in the field.

(c) To explain the inconsistent results obtained on certain main radar equipment.

When "calibration" is referred to without further qualification in this Instruction, the calibration of output amplitude is intended; frequency calibration is not in question.

Uses of a signal generator

2. A signal generator fulfils two main functions: the alignment of a receiver, or the measurement of its bandwidth or sensitivity. Alignment commonly consists of adjusting the tuned circuits of one or more amplifying stages until maximum response is obtained, and for this purpose an unmonitored signal source of stable amplitude suffices; the accuracy or otherwise of the calibration of a signal generator used for this purpose is therefore of no importance, and alignment is therefore not relevant to the considerations of this Instruction.

3. The **measurements** most commonly made with a signal generator are three, viz:—

(a) bandwidth,

(b) stage-by-stage sensitivity

(c) overall sensitivity

These will be considered in turn.

4. The measurement of **bandwidth** typically involves measurement of receiver output at two frequencies relative to that at a third frequency intermediate between them. Apart from these variables the conditions of measurement remain unchanged: no absolute measurement is called for, and the sole requirement is that the calibration of the signal generator, in whatever units it may have been expressed, should remain self-consistent over the ranges of output and frequency concerned.

5. **Stage-by-stage** sensitivity measurements are of use in locating a faulty stage of amplification, the specification of the receiver under test being known in terms of the output at a given point which should be obtained from a known input to a given stage. The precise form in which the specification is given is not important, but it is important to note that such figures must be reproducible for **any** specimen of the given type of receiver in conjunction with **any** specimen of appropriate signal generator. This requirement can only be fulfilled if the signal generator calibration can be referred to some external

standard, so that the results obtained are a true measure of the receiver performance, and do not depend on the specimen or type of signal generator employed. Finally, it is relevant to note that these measurements are typically taken with the signal generator output applied across grid-cathode of the stage under test; the impedance presented to the signal generator by receiver grid-cathode circuits other than the input (see para. 6) is commonly, though not invariably, high compared with that of the signal generator output, and predominantly resistive.

6. **Overall** sensitivity measurement is intended to provide a measure of the response of a receiver to a signal such as it might receive in operation, and is no more than a special case of the process described in para. 5. The reason for treating it separately is that it is uncommon to take this measurement across grid and cathode of the first R.F. stage (if any); the test signal will normally be applied to the aerial input socket or its equivalent. This may lead to an input transformer, a mixer box, or some similar device, whose characteristic is that it is normally accustomed to receive a signal at a particular low impedance, the circuit being designed to extract maximum power at that impedance. It is also important to note that overall sensitivity may be measured in terms of **either** receiver output for a given input (i.e. overall gain) **or** the input required to produce a specified signal/noise ratio. Of these alternatives, the latter provides by far the more valuable and reliable index of the performance of a receiver, particularly of a receiver intended to operate at ultra-high frequencies; signal/noise measurement is therefore frequently called for, and it is relevant to observe that its measurement calls for appreciably greater care to obtain accuracy than does that of overall gain.

The attenuator and its calibration

7. Two types of attenuator are in common use: the "piston attenuator", whose special properties have been described in D.M.E.(I.) T.I. Desc. Tels. G. 03, dated 22nd May 1944 and the "slide-wire", which may be regarded as a form of autotransformer, normally used in conjunction with a decade-box (a resistance attenuator permitting selection in powers of ten). These will be considered in turn.

8. It is vitally important that the impedance presented to the attenuator tube by the collector mechanism be completely unaffected by the presence or absence of the receiver to which the output cable is to be connected. Failure to observe this requirement will in general result in reflection from the collecting mechanism, which will invalidate the attenuator calibration. This difficulty is overcome by the provision of an output cable, without which the signal

generator may not be used, of 10 or 20 db. attenuation and of suitable impedance, effectively buffering the signal generator from the receiver under test. Now, a known power is being injected into the attenuator, and a measured fraction of this power is available for application to the receiver. This is not, it should be noted, a measured **voltage**; for though it may be valid to measure the voltage across the collector, or at the end of the output cable, it does not follow that this voltage is available for application to the receiver.

9. Similar considerations apply to the slide-wire attenuator. Although it may be argued that auto-transformer devices are primarily voltage transformers, the signal generator is so designed that the slide-wire is, as far as practicable, presented with a constant impedance, in order that its calibration may not be affected by alterations in impedance at the receiver: the voltage-fraction produced is therefore available at a given fixed impedance, and a certain definite **power** is represented by any slide-wire setting.

10. Voltage measurements at very high frequencies (i.e. approaching the centimetric region) are neither convenient to make nor easy to interpret. Signal generators operating at such frequencies are therefore normally calibrated direct in power (i.e. in watts, not in volts). Furthermore, as stated in para. 6, measurement of signal/noise ratio is of paramount importance at these frequencies, and the quantitative estimation of noise, for reasons which need not be considered here in detail, is essentially a **power** measurement. However, there are in Service use signal generators operating at 300 Mc/s and below, which are calibrated in volts, and it is important to appreciate the precise implications of such calibration. One such signal generator in common use has a 30 Ω output cable; the calibration represents the voltage read at the end of this cable when properly terminated by a 30 Ω resistive impedance. If the cable is terminated in any other way, the calibration is invalid. This, or a similar condition, is always implicit in measurements taken with a voltage calibrated signal generator.

VOLTAGE CALIBRATION IN PRACTICE: TYPICAL CASES

I.F. high impedance grid-cathode input

11. Still considering the 30 Ω output, suppose the proper terminating resistor to be in place, and its two ends connected to grid and cathode of an I.F. stage operating below 100 Mc/s. (Both signal generator output and the I.F. input are assumed unbalanced, one side of each being at earth potential). The input impedance of such a stage is commonly very high compared with 30 Ω and predominantly resistive, the stage drawing practically no power from the input. It is clear that the connection of a high resistance across the terminating impedance will not materially affect the voltage developed across it; provided the correct terminating unit is in use, the voltage as indicated by the signal generator will be a sufficiently accurate measure of the voltage developed between grid and cathode of the stage under

test. Voltage calibration is clearly suitable for measurements of this type.

R.F. low impedance input

12. Assume that the aerial input socket of a receiver has normally an input impedance of 80 with one side at R.F. earth. If the 30 Ω cable is directly connected to this input, without its terminating impedance, the condition for validity of calibration will be unfulfilled, and the calibration will be meaningless. If the 30 Ω terminating unit is in circuit, the 80 Ω shunt will extract appreciable power, and the calibration will again be invalidated. It might be argued that the calibration, though meaningless, would be repeatable and therefore useful as a comparative measure. This is not true, since there will inevitably be small variations in input impedance between different receivers; since this impedance is of the same order as that of the terminating unit, these variations will appear as unknown errors in calibration, and therefore in sensitivity measurement.

13. The concept of "matching in" must now be examined. It is conceivable, though most improbable, that adjustments might be provided which would enable the normally 80 Ω input to be matched to the 30 Ω cable, in which case the calibration would be valid. However, the sensitivity measurement thereby obtained would **not** be valid; for the input is accustomed to receive a signal at 80 Ω and its sensitivity at 30 Ω would provide no adequate measure of its sensitivity under normal operating conditions. This results in the phenomenon of the receiver which appears sensitive when on bench-test, but insensitive in its operational equipment. If, on the other hand, an impedance transformer is employed as a matching unit, the voltage across the 80 Ω side is not that across the 30 Ω side: assuming that some slight matching adjustment has been necessary, the voltage may not be calculable. However, the **power** can be transferred with a minimum of loss, and although the voltage calibration may be meaningless, power calibration would remain correct and significant. Furthermore, at conditions near the matching-point an appreciable degree of mismatch will result in only a relatively small change in transmitted power, but in a relatively large change in measured voltage. (It is, incidentally, worthy of note that an input matched for maximum transference of power is not matched for maximum signal-to-noise ratio, since the maximum do not coincide). It is clear that power calibration is desirable for measurements of this type.

VOLTAGE CALIBRATION IN PRACTICE: ANOMALOUS CASES

Low impedance grid-cathode inputs

14. Three cases will be met in practice:—

(a) Input to a pentode R.F. stage operating in the region of 200 Mc/s.

(b) Input to a grounded-grid triode R.F. stage in the same region.

(c) Input to the first I.F. valve after a head-amplifier, in cases where the I.F. signal is transmitted via a low impedance coaxial cable.

These cases will not be well served by either method of calibration; the impedance is too low to allow the voltage calibration to be regarded as trustworthy, and the existence of suitable matching arrangements for power calibration is neither probable nor desirable. The method of calibration of the signal generator is therefore of no great importance in these examples, since it must in any case be regarded as unreliable.

A.A. No. 1, Mk. II, receiver

15. This equipment presents several peculiar difficulties, and the results of sensitivity measurements obtained with different signal generators on the same receiver, or with the same signal generator on different receivers, have been so inconsistent that an explanation of the difficulties is desirable.

16. The input impedance is nominally $80\ \Omega$ resistive and balanced. If it be assumed for the moment that this is actually the case, it is clear that the impedance is sufficiently low to warrant power calibration. Now, no other radar receiver in extensive use operates at so low a frequency, and no other has a balanced input; but, provided the characteristics of the receiver were as stated, an alternative calibration in power could be provided for signal generators operating in this range. However, the receiver is of very early design, and at the time of its design the technique of accurate impedance measurements at these frequencies was itself in a fairly early stage, nor was its importance so fully realised as it is now. In consequence, it has been established that the input impedance may vary widely as between receivers, a value as high as $220\ \Omega$ having been reported; this impedance may have a large reactive component, and the balance is neither accurate nor consistent in its inaccuracy. Now, the receiver is not provided with matching arrangements that could eliminate these variations, and, if the good matching essential to the validity of power-calibration measurements is to be attained, a special matching unit, capable of matching this widely varying impedance into, say, a $30\ \Omega$ line would be required, and it would have to operate over the entire frequency-band used by this equipment. Such a unit would be complex, difficult to design and cumbersome to operate; it is generally agreed that the design and production of such a unit, and therefore the power calibration of the relevant signal generator, would be technically unprofitable.

17. The most practical compromise, therefore appears to be the following:—

(a) To design a simple unbalance-to-balance terminating unit for the signal generator, designed to match into the theoretical $80\ \Omega$ resistive balanced input. The design of even such a unit to be complete frequency-insensitive over the whole band is a matter of considerable difficulty; but since this is in this case a comparatively minor inaccuracy, the most important provision is that all terminating units in use shall be as near identical in characteristics as practicable.

(b) To retain the existing voltage calibration, since power calibration is not justified if reasonably good matching cannot be obtained.

(c) To realise that specification figures for the sensitivity of this receiver can be at least only a guide to the order of result that may be expected. Except in unusually favourable circumstances, specification figures will be unrepeatable with changes of either signal generator or receiver.

Conclusions

18. From the above we may draw the following conclusions:—

(a) For low-frequency high-impedance grid-cathode inputs, voltage calibration is preferable.

(b) For radio-frequency overall sensitivity measurements, power calibration is preferable.

19. Apart from the anomalous cases already discussed—in all of which it has been noted that the calibration used is immaterial, since both types are untrustworthy—a convenient distinction can immediately be made on frequency, and a dividing line fixed, such that voltage calibration is used at frequencies below this line and power calibration above it. No rigid decision need be made, but it may be assumed that the dividing line will be in the neighbourhood of 100 to 150 Mc/s.

TECHNIQUE OF CALIBRATION

Power calibration

20. It will be clear that, even though the signal generator attenuator were a piston of known linearity and established accuracy, at least one direct R.F. power measurement must be made. The most convenient devices to effect this measurement are the bolometer lamp, the thermistor, and the thermocouple. Bolometer lamps are, for various reasons, unsuitable for this purpose; insufficient experience has yet been gained in the use of thermistors to permit of their general adoption; and the thermocouple, whose characteristics are well known, is therefore used in all signal generator calibration equipment in development or production.

21. A difficulty which arises immediately is the fact that a thermocouple cannot be used for accurate power measurement at levels below about 1 mW, whereas R.F. signal generators at present in use work at levels more nearly approaching 1 W. Some form of transfer instrument is therefore required.

22. Reference should now be made to fig. 1. The calibration of a signal generator involves the following operations:—

(a) The signal generator, its attenuator set to the power level at which it is desired to take the measurement, is connected through a matching section to a standard comparison receiver. The gain of this receiver is set to give a convenient deflection on a meter reading second detector current.

(b) The signal generator is now removed, and replaced by a reference oscillator, whose characteristics are its capability of delivering up to several milliwatts, and the established accuracy of its piston attenuator. The gain of the receiver is left unchanged from operation (a) and the reference oscillator

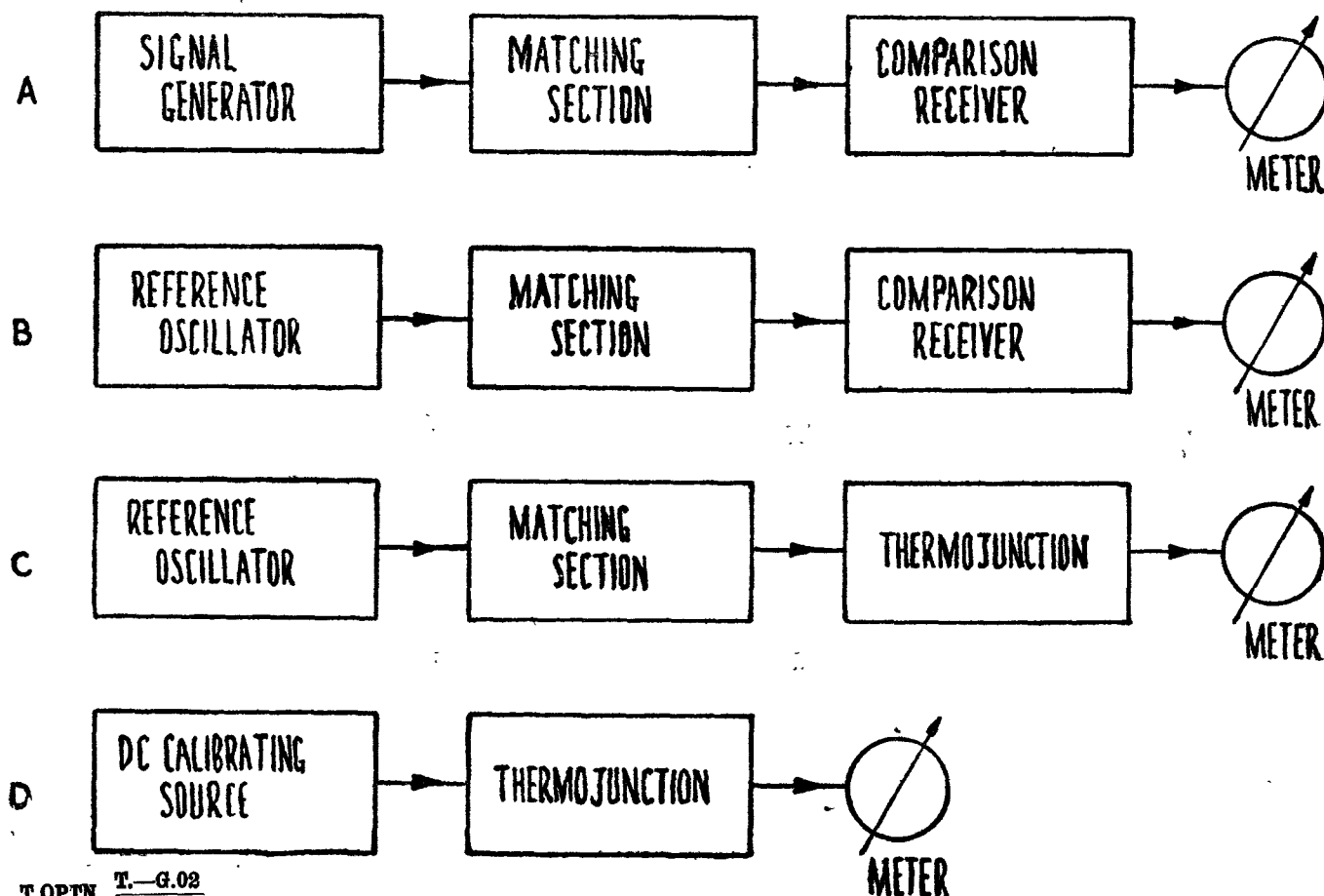


Fig. 1 Block diagram showing operations involved

adjusted until the meter reading is the same as before, re-matching if necessary. The reference oscillator is now delivering the same power as was the signal generator.

(c) The comparison receiver is now removed, and the reference oscillator is connected through the matching section direct to a thermocouple, whose output is metered; the reference oscillator attenuator is adjusted until a convenient meter deflection is obtained.

(d) The reference oscillator and matching section are replaced by a D.C. calibrating source whose amplitude is adjusted until the thermocouple output is the same as that observed in (c). The voltage across, and the current through, the thermocouple filament are measured.

23. It will be clear that operation d) permits the power dissipated in the thermocouple filament to be calculated by simple multiplication, and this power will be the same as that dissipated in the filament in operation (c). The matching loss of the system is known, and so the power output of the reference oscillator in operation (c) is known. From a comparison of the reference oscillator attenuator readings in operations (c) and (b), the power delivered by the reference oscillator in operation (b) is known, and

therefore the power delivered by the signal generator in operation (a) is known. By taking such measurements at different settings of the signal generator attenuator, the attenuator can be calibrated at as many points as desirable. It is important to note that the validity of the process is completely dependent on the accuracy of the attenuator of the reference oscillator, which must therefore be manufactured with considerable accuracy and treated with due care in operation. It would, in theory, be possible to calibrate the reference oscillator directly in terms of power, but, since the oscillator output will inevitably fluctuate to some extent, a direct thermocouple measurement is made for each calibration measurement. Errors due to fluctuations of thermocouple characteristics are avoided by operation (d), which is carried out within one or two minutes of operation (c).

24. A complete calibrator will be seen to consist of the following items:—

- (a) Reference oscillator covering the appropriate frequency-band.
- (b) Comparison receiver capable of operating over this band.
- (c) Matching section.
- (d) Thermocouple, with D.C. calibration facilities.

D. M. E. (INDIA)
TECHNICAL INSTRUCTIONS

OPERATION
TELS. G. 02

(e) Suitable meter, to be switched into relevant circuit as required.

Voltage calibration

25. Voltage calibrators are in essentials similar to power calibrators, except that the thermocouple is replaced by a valve voltmeter.

Alternative techniques

26. The above system is the most suitable where, as in signal generators intended for use with Army

radar equipments, only a relatively narrow frequency-band is required. If a very wide frequency range has to be covered, an instrument of this type becomes inconveniently complex, and an alternative technique is used. This will not, however, be required for Army use, and all calibrators now in development or production will be of the type described above. Detailed descriptions of the calibrators will be the subject of separate Instructions as the equipment becomes available.

END

(3521/7/MG/ME-11)

D. M. E. (INDIA)
TECHNICAL INSTRUCTIONS

DESCRIPTION TELS. G. 04

WAVE-GUIDES
OUTLINE OF PRINCIPLES

RESTRICTED

**THE INFORMATION GIVEN IN THIS DOCUMENT
IS NOT TO BE COMMUNICATED, EITHER DIRECTLY
OR INDIRECTLY, TO THE PRESS OR TO ANY PERSON
NOT AUTHORISED TO RECEIVE IT.**

Issue 1, 24th May 1944

WAVE-GUIDES

Outline of principles

[Based on E. M. E. R. Tels. A. 011, Issue 1]

CONDITIONS UNDER WHICH ELECTROMAGNETIC WAVES EXIST NEAR CONDUCTORS

1. A plane electromagnetic wave in free space has oscillations of the same phase and amplitude at all points in its wave-front (that is, in a plane perpendicular to the direction of propagation of the wave). If a perfect conductor is in the presence of an electromagnetic wave, the necessary conditions at the surface of the conductor are that both the component of the electric vector parallel to the surface $E_{\parallel}$ and the component of the magnetic vector perpendicular to the surface ($H_{\perp}$) shall be zero. In order that this may be so, the amplitude of oscillation must in general be non-uniform along the wave-front. A single plane wave,

therefore, cannot as a rule exist near a conductor, but by a combination of waves, it is possible to satisfy the conditions at the conductor.

A PLANE WAVE INCIDENT UPON A PLANE CONDUCTOR

The Electric Vector

2. Let there be incident upon a plane perfectly conducting surface a plane wave having its electric vector parallel to this surface. The incident wave is reflected with a reversal of phase of the electric vector so that the resultant at the surface may satisfy the condition $E_{\parallel} = 0$. Also, the incident and reflected wave-fronts make equal angles with the conductor. If

VARIATION OF AMPLITUDE IN A LINE PERPENDICULAR TO THE SURFACE

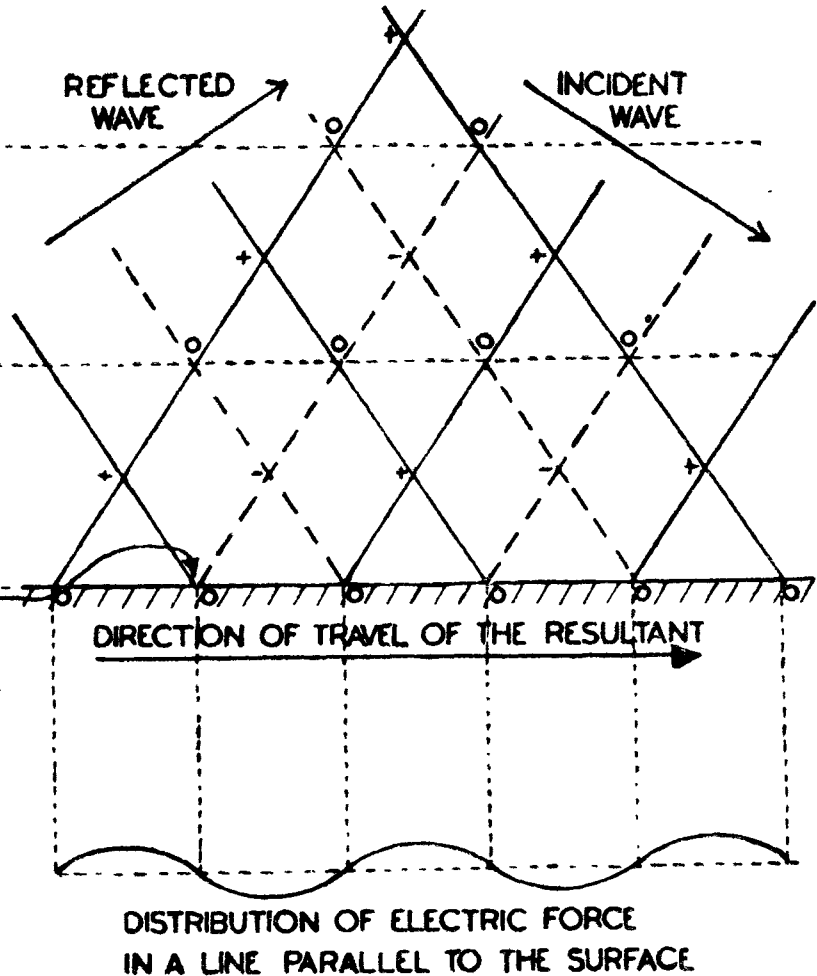
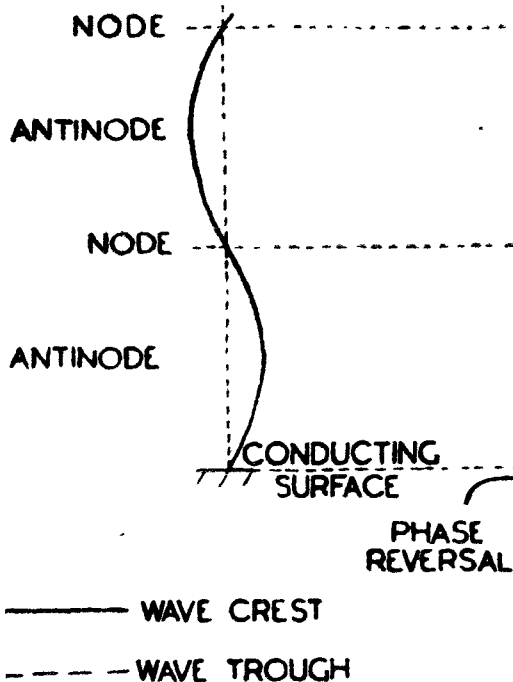


Fig. 1

DESC. T-C 04.
 1-1

the incident wave-front is parallel to the surface, combination of incident and reflected waves gives a resultant which is a pure standing wave along the normal to the surface. If the incident wave-front is not parallel to the surface, combination of incident and reflected waves gives a resultant, travelling parallel to the surface, which has a sinusoidal variation of amplitude along the normal. This is illustrated in figure 1, which shows how the incident and reflected wave-fronts combine.

3. At points marked O the incident and reflected waves are out of phase and the electric fields cancel, these points therefore lie in nodal planes which are parallel to, and spaced at equal intervals from, the reflecting surface. At the points marked + and - the electric fields are in phase and reinforce each other, these points therefore lie in planes of maximum amplitude midway between the nodal planes. Variations of amplitude between nodal planes are indicated on the left of the figure and the distribution of electric force along the direction of propagation of the resultant wave is shown below the figure. The amplitude of oscillation between nodal planes varies sinusoidally and there is a 180° phase change on crossing a nodal plane. Since there can be no transfer of energy across a nodal plane the energy associated with a particular sinusoidal loop of amplitude between nodal planes is completely independent of any other, the energy can therefore be considered as travelling down channels between nodal planes.

4. The electric vector is zero at a nodal plane and a further reflecting surface can be placed there without disturbing the electric field. Also at a plane perpendicular to the electric vector, $E_{\parallel}$ and $H_{\perp}$ are

always zero, so that reflecting surfaces can be placed parallel to this plane without disturbing either the electric or magnetic fields and since the wave is travelling parallel to this plane, there is no transfer of energy across it. It remains now to show that the magnetic field is also unaffected by placing a conductor at a nodal plane.

The Magnetic Vector

5. In figure 2 are shown the magnetic vectors of the same wave system as figure 1. The H vectors of the incident and reflected waves lie in the planes of the wave-fronts and are perpendicular to the E vectors, and hence they must lie in the plane of the figure.

6. This figure shows the manner in which the incident and reflected H vectors combine. At both the reflecting and the nodal surfaces, the component of H perpendicular to the surface is zero since the angles of incidence and reflection are equal. Hence the boundary condition for $H_{\perp}$ is satisfied at the reflecting surface itself, and also at a reflecting surface placed at a nodal plane. There is a component of H in the direction of propagation which is a maximum at the nodal planes, and zero mid-way between.

7. From paras. 2 and 3, therefore, it can be seen that electromagnetic waves will travel down a hollow conductor of rectangular cross-section, since the boundary conditions for $E_{\parallel}$ and $H_{\perp}$ at the inner surface of the conductor can all be satisfied. The oscillations will be completely independent of any outside phenomena since it has been shown that there is no transfer of energy across surfaces at which these boundary conditions are satisfied. There will be a component of H in the direction of propagation and if a cross-section

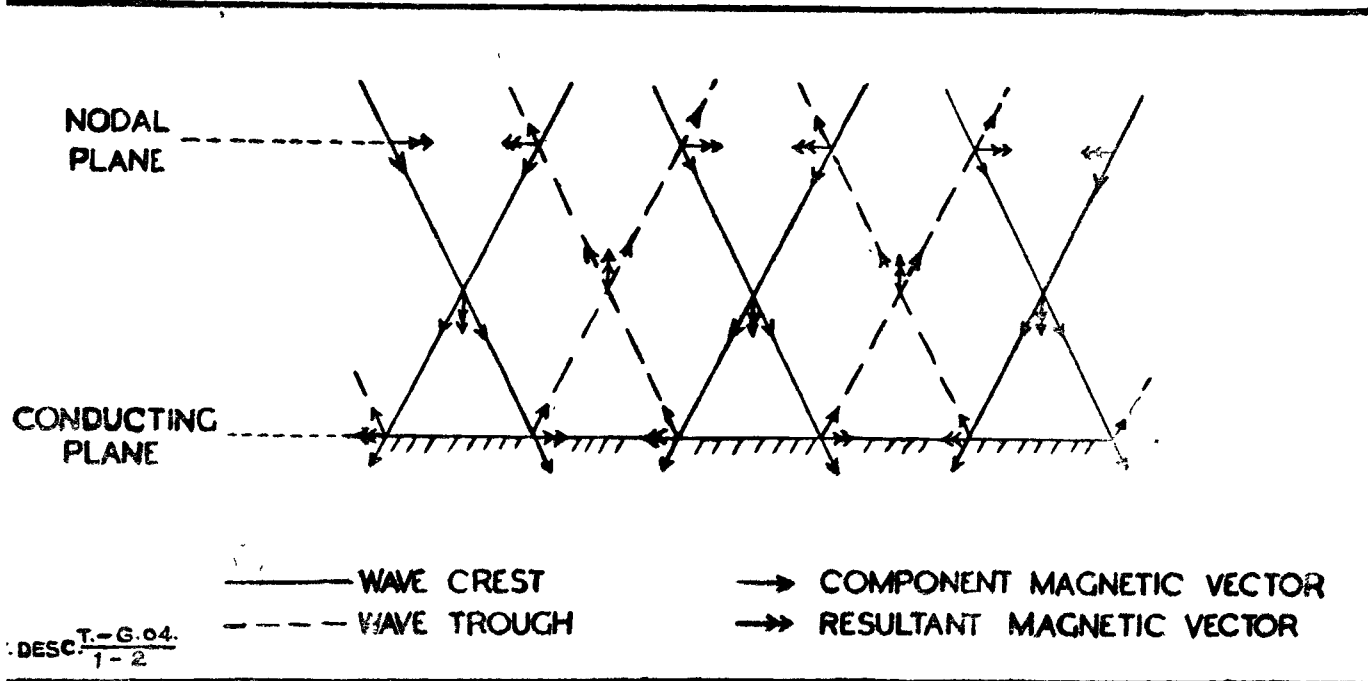
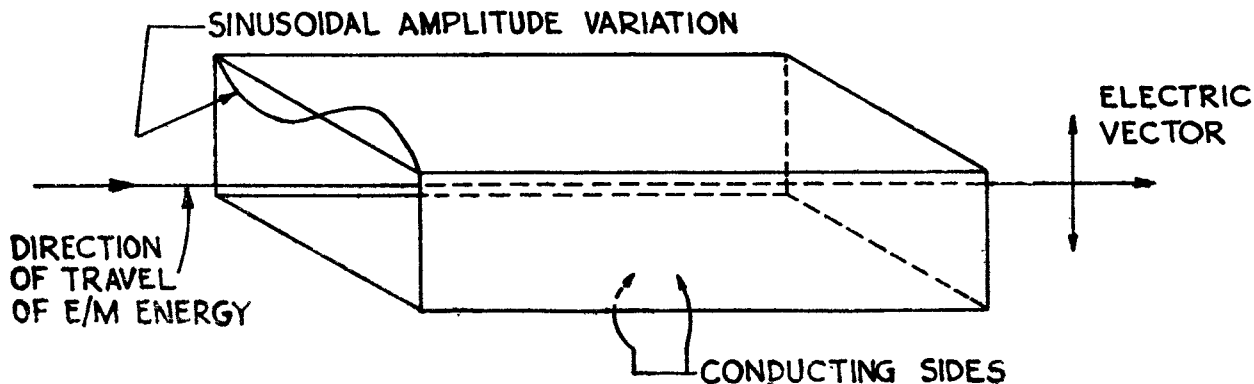
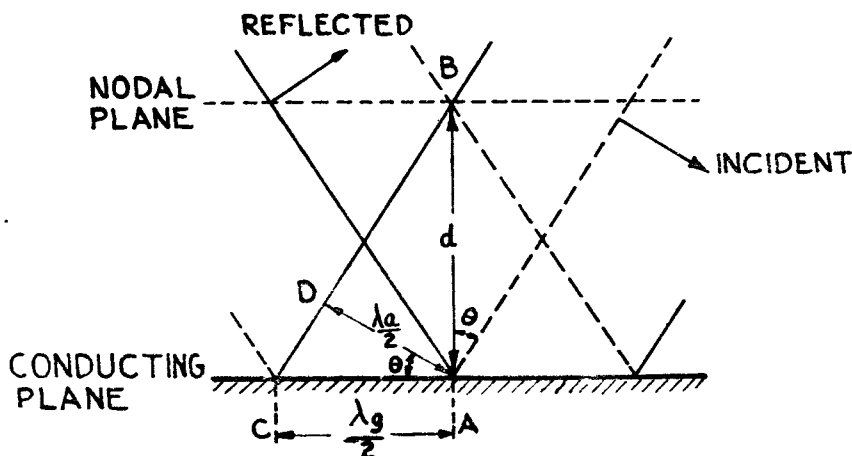


Fig. 2



T.DESC. T-G.04
1-3

FIG. 3



T.DESC. T-G.04
1-4

FIG. 4.

of the conductor is taken, there will be a sinusoidal variation of amplitude of E in one direction, as shown in figure 3. Transmission systems making use of this principle are known as wave guides.

MATHEMATICAL ASPECTS

8. In figure 4:—

A=point of incidence of a particular wave trough.

B=point of intersection of the preceding reflected trough and the following incident crest.

C=point of incidence of the following crest.

CD=projection of CA upon CB.

λ_g =wave-length of the resultant wave in the guide.

λ_a =wave-length of the component plane waves (i.e. the wave-length of the radiation in free space).

θ =angle of incidence.

Then from $\triangle s=ADC$ and ADB

$$AD=AC \cos \theta$$

$$=AB \sin \theta$$

$$\text{i. e. } \lambda_a = \lambda_g \cos \theta$$

$$=2d \sin \theta$$

Eliminating θ :—

$$\left(\frac{\lambda_a}{2d}\right)^2 + \left(\frac{\lambda_a}{\lambda_g}\right)^2 = 1$$

$$\lambda_g = \frac{\lambda_a}{\sqrt{1 - \left(\frac{\lambda_a}{2d}\right)^2}} \dots \dots \dots (1)$$

9. In equation (1), there are three different cases to be considered

$$(a) \frac{\lambda_a}{2} < d$$

Equation (1) gives a real, finite solution for λ_g . The

wavelength as measured in the wave-guide is greater than the wave-length in free space. Also, since the frequency is constant, the velocity in the wave-guide is therefore greater than the velocity in free space. This appears to conflict with relativity theory which states that the greatest possible velocity is the velocity (c) of electromagnetic waves in free space. It is not the case, however, since the wavelength in the guide is related to phase velocity (v_p) and it is group velocity (v_g) which is limited. The condition is that

$$v_p \cdot v_g = c^2 \dots \dots \dots (2)$$

The phase velocity is the velocity of propagation of an actual phase of the wave-motion down the wave-guide and is of value $\frac{c}{\cos \theta}$. The group velocity is the velocity of propagation of the electromagnetic energy itself down the wave-guide and is therefore the velocity of the component plane waves resolved along the wave-guide, this is $c \cdot \cos \theta$. Thus equation (2) is satisfied.

$$(b) \frac{\lambda_a}{2} > d$$

Equation (1) gives an imaginary value for λ_g and therefore the theory previous to it does not hold in this case. A satisfactory analysis can be obtained mathematically, beginning with Maxwell's fundamental electromagnetic equations. The results only will be stated here. This type of wave-guide is known as an evanescent wave-guide, since it is found that if electromagnetic energy is fed into it the phases of the oscillations are independent of the position, and the amplitude decays exponentially along the guide. This is a peculiar type of standing exponential oscillation which also occurs in the theory of the penetration of electromagnetic waves into an imperfectly conducting material.

$$(c) \frac{\lambda_a}{2} = d$$

This is the transition condition between cases (a) and (b). Equation (1) gives an infinite wave-length and equation (2) shows that the phase velocity is infinite but the group velocity is zero. Therefore, oscillations at all points down the wave-guide are in phase, equal in amplitude, and no energy is transmitted down the wave-guide.

MODES OF OSCILLATION IN A RECTANGULAR WAVE-GUIDE

10. In para. 2, the incidence of an electromagnetic wave upon a conductor was considered, in which the particular assumption was made that the E-vector was parallel to the conductor; one of the conclusions from this was that the H-vector has a component in the direction of propagation of the resultant wave. If the H-vector had been taken as parallel to the conductor instead of the E-vector, similar results would have been obtained, except that the E-vector would have had a component in the direction of the propagation of the resultant. These two observations may be summarised in the following rule that if one vector of an electromagnetic wave has a non-uniform amplitude over the wave-front, then the other vector must have a component in the direction of propagation of the wave. The types of wave which may be transmitted

down a wave-guide can therefore be divided into two groups :—

(a) those which have a component of H in the direction of propagation. These are known as *H-waves*.

(b) those which have a component of E in the direction of propagation. These are known as *E-waves*.

11. Each wave in these groups can be distinguished from other members of its group by referring to the variation in amplitude of the wave over the cross-section of the guide. These are called *geometrical modes* of oscillation. In para. 2, it was shown that the cross-section must contain a whole number of sinusoidal loops of amplitude. Moreover, the amplitude may be made to vary simultaneously in two directions, namely, parallel to the adjacent sides of the wave-guide cross-section.

12. The following nomenclature is therefore used for the modes of oscillation in a rectangular wave-guide :—

An $H_{m,n}$ wave is defined as an H wave having m sinusoidal loops of amplitude in one direction and n loops in the direction at right-angles to this. An $E_{m,n}$ wave is similarly defined. For example, figure 5 indicates the variation of amplitude over the cross-section of a wave-guide carrying an $H_{1,0}$ wave.

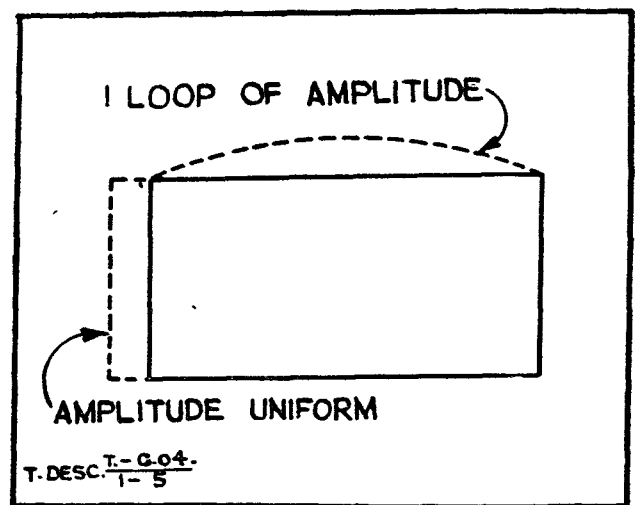


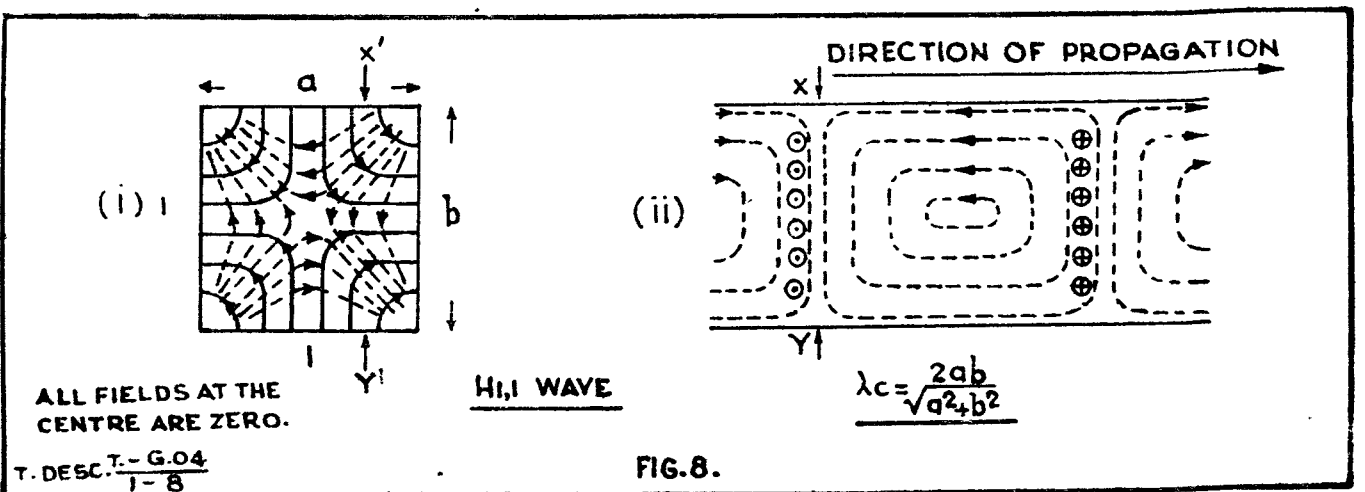
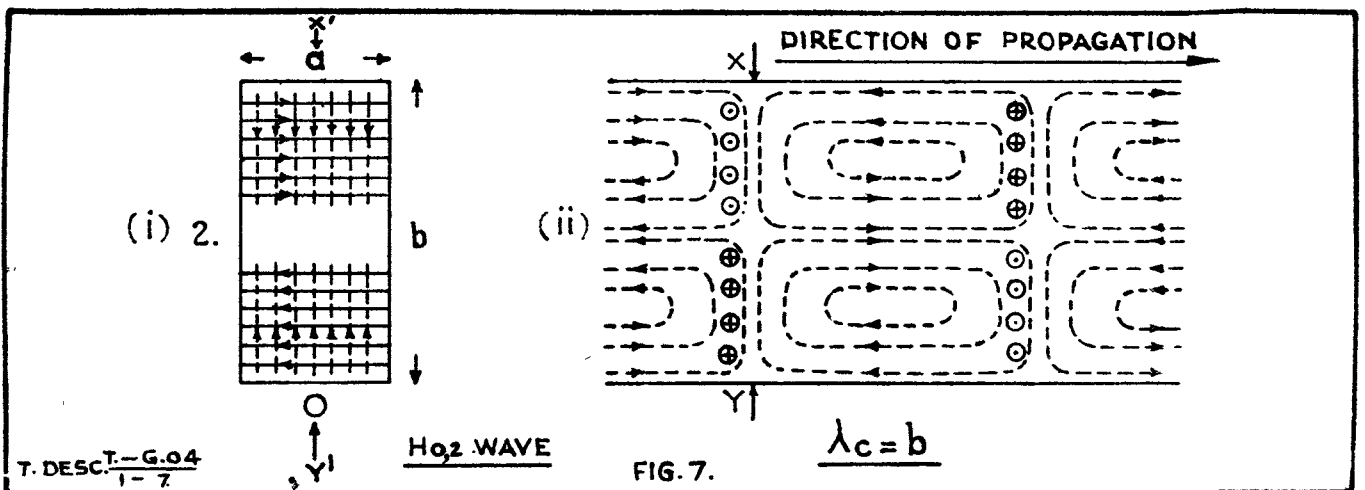
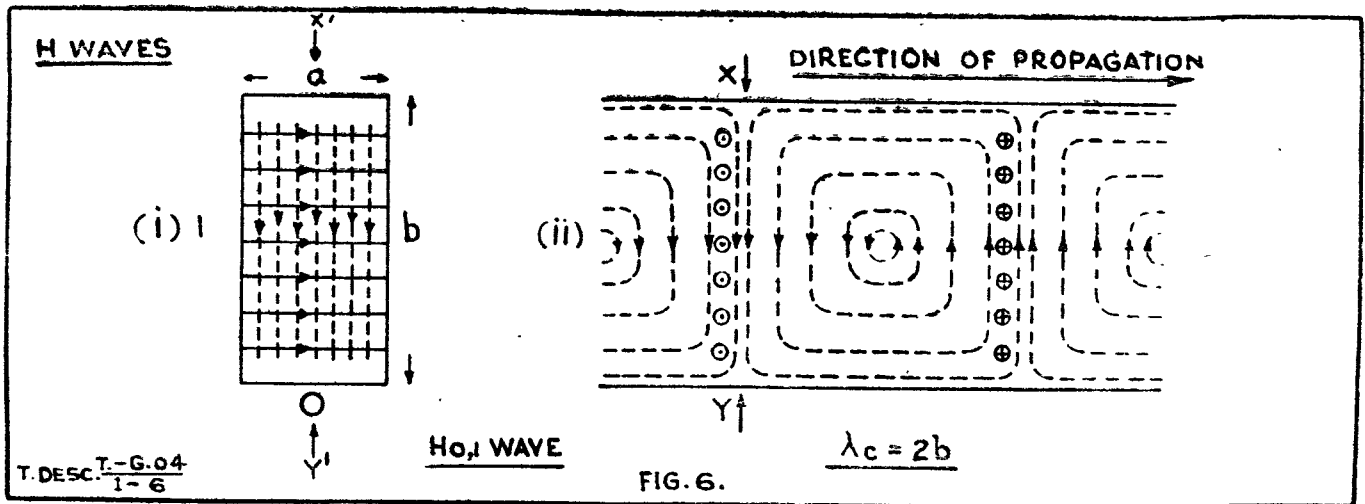
Fig. 5

13. A knowledge of the field configuration of various modes of oscillation is of practical importance and the more important of these will be drawn in the following figures. In each case, a cross-section and a longitudinal section of the wave-guide are illustrated. It is found that the modes of oscillation of a wave-guide are like those of a vibrating membrane whose shape is similar to the cross-section of the wave-guide.

Key to Figures 6—13

14. (a) Representation of fields in the plane of the figure.

- > E lines of force
- > H lines of force



Discontinuity in the line of force indicates that the field no longer has a component in the plane of the figure.

(b) Representation of fields normal to the plane of the figure.

⊙ E line passing normally out of the figure

⊕ E line passing normally out of the figure

⊙ H line passing normally out of the figure

⊕ H line passing normally into the figure.

15. No attempt has been made in the figures to indicate magnitudes of vectors ; the lines of force drawn on each diagram are envelopes of the particular vector under consideration. In the case of rectangular wave-guides, an indication of the number of sinusoidal loops of amplitude along each side is given by placing appropriate numerals near the corresponding side of the cross-section.

λ_c = Critical wavelength given by equation :—
 $c = \lambda_c \eta_c \dots \dots \dots (3)$

Where η_c is the lowest frequency at which the mode passes freely down the wave-guide.

16. No E wave involving a zero suffix can exist in a wave-guide of rectangular cross-section. The zero suffix implies no change of amplitude in one direction and as the component of E in the direction of propagation must be zero at the sides, it must be zero everywhere.

$H_{1,0}$ is an exception to this, however, and for a given input frequency, the smallest cross-sectional dimensions are necessary for propagation. This type of wave can therefore be obtained in a very pure state by choosing the values of a and b such that all modes other than $H_{1,0}$ are attenuated evanescently. As a result of this $H_{1,0}$ waves are used most commonly, a secondary advantage is that the wave-guide is in its most compact form.

CIRCULAR WAVE-GUIDES

17. From figures 10 to 13, it is evident that, of the four types of wave illustrated, H requires the smallest diameter for propagation at a fixed frequency ; more generally, it can be stated that in a circular wave-guide no other mode requires such a small diameter for propagation at a fixed frequency. H_1 waves are therefore widely used since they are easily produced in the pure state by making the diameter of such a size as to exclude all but H waves, this being also the most convenient size for practical use.

18. Figures 12 and 13 also show that the electro-magnetic field patterns of E_0 and E_1 waves are very similar to those of concentric and shielded twin cables respectively. The currents carried by the inner conductors of the cables have, however, been replaced by electric displacement currents in the case of the corresponding wave-guides.

19. The establishment of the modes of oscillation of a wave-guide of circular cross-section is similar to,

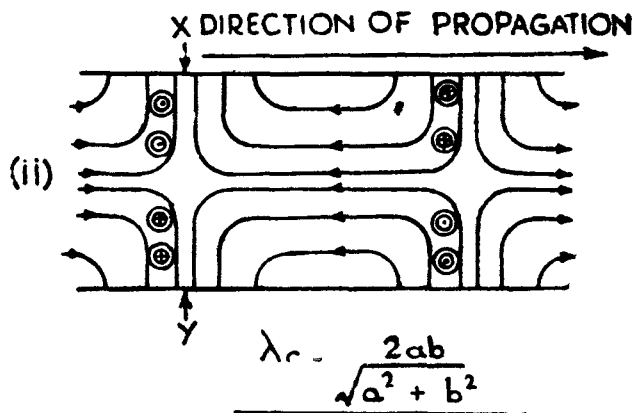
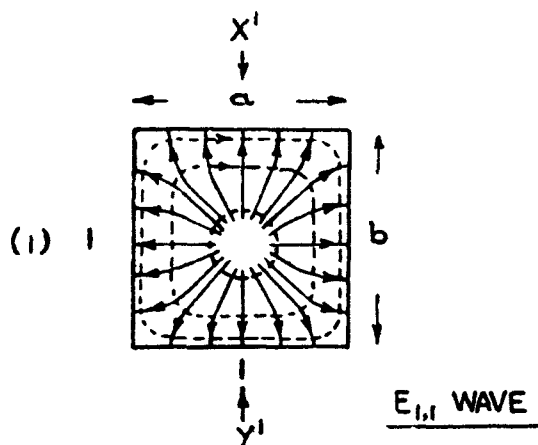


FIG. 9.

T DESC. T. G 04
1-9

From the values of λ_c given in the figures, it is evident that when the condition for propagation of one mode is fulfilled, then others may hold automatically. Thus, for example, the conditions for propagation of $H_{1,1}$, $H_{1,0}$ and $H_{0,1}$ waves, may be automatically satisfied by a wave-guide which will freely pass an $E_{1,1}$ wave. $H_{1,1}$, $H_{1,0}$ and $H_{0,1}$ modes may therefore be expected as impurities in a wave-guide down which $E_{1,1}$ is passing.

but mathematically more difficult than that for a rectangular wave-guide, and involves the use of Bessel Functions. The modes can be divided into E and H waves as in rectangular wave-guides, but the geometrical modes are distinguished by the order of the Bessel Function used in their derivation. For example, an E_0 wave is an E wave involving the Bessel Function of order zero. H waves are similarly classified. The

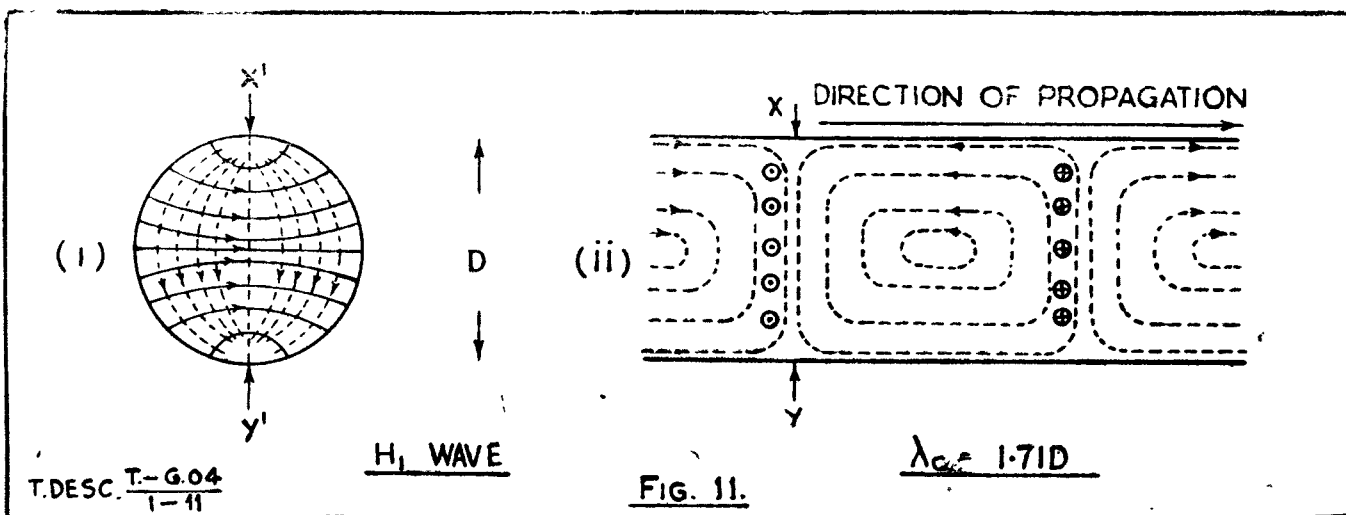
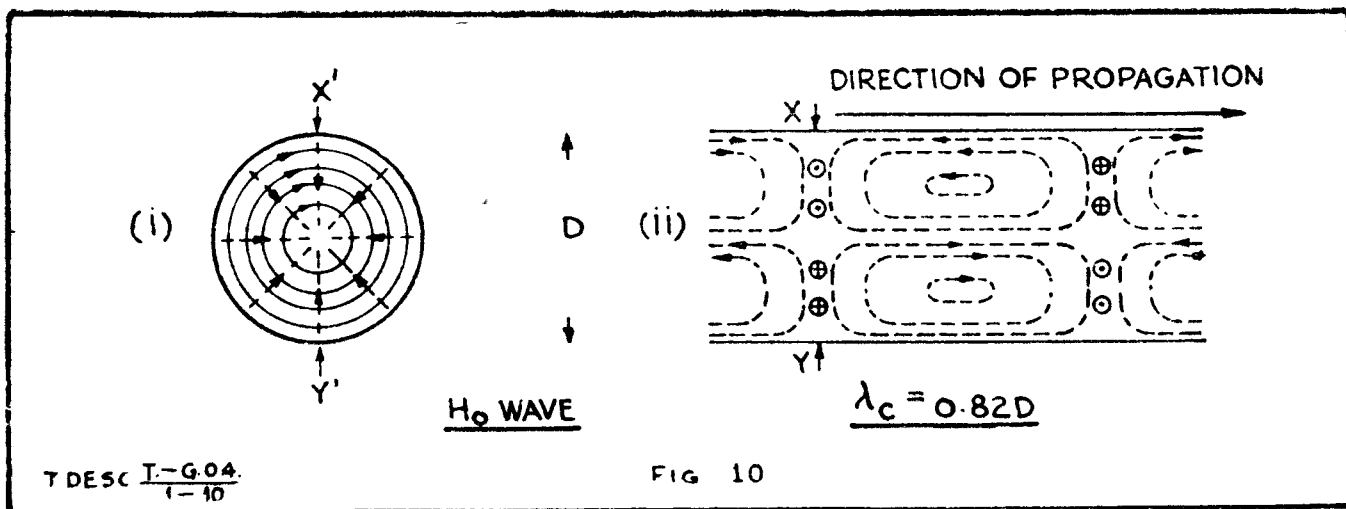
modes of oscillation of circular wave-guides are analogous to the modes of vibration of a circular membrane. Also, there is a limiting wave-guide radius below which propagation is impossible at a given wavelength.

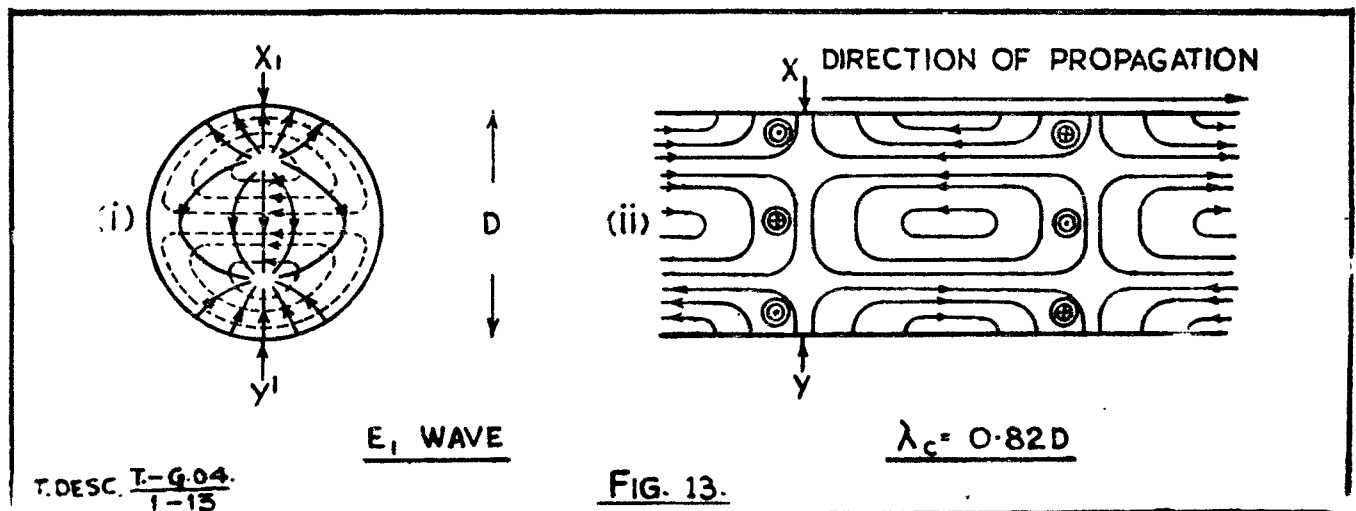
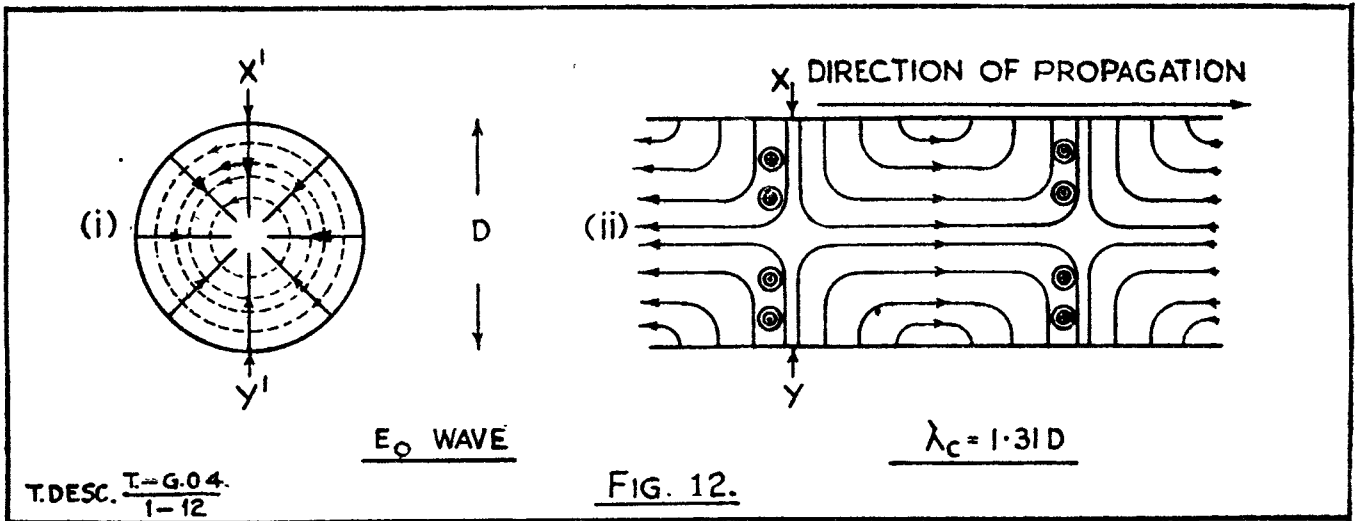
EXCEPTIONAL TYPES OF WAVE-GUIDES

Wave-guides of relatively small dimensions

20. At a wavelength of the order of 10 cms. ordinary

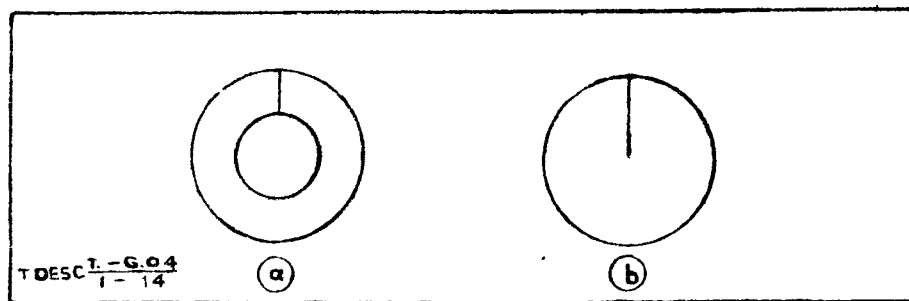
circular wave-guides are somewhat cumbersome for practical use owing to their large diameters. One method of overcoming this is to fill the wave-guide with dielectric. Since λ_c is proportional to the square root of the dielectric constant of the medium, the wave-guide diameter may be correspondingly reduced. This leads to a sharp rise in attenuation (see para. 24) and defeats the main object for which wave-guides are used.





21. Another method of reducing the dimensions is to "fold" back the sides of a rectangular wave-guide till they are coincident, and thus form the cross-sectional shape as shown in figure 14 (a). This is known as a septate wave-guide; it has the appearance of a concentric transmission line, the inner of which is

connected to the outer by a plane conducting septum. If now the inner is reduced to zero diameter, the septum alone is left. This is known as a "degenerate septate" and is shown in figure 14 (b). The latter type of wave-guide is easier to construct and therefore more commonly used.



(a) Septate.

(b) Degenerate septate.

Fig. 14. Wave-guide cross-sections

22. In figure 6, it was seen that $\lambda_c = 2b$ for an $H_{1,0}$ wave and this is the largest wavelength which can be passed down the wave-guide; if the wave-guide were now folded into a septate, it would be expected that the longest wavelength which could be passed down the wave-guide would be approximately twice the mean circumference of the two cylinders. This is actually the case and λ_c is therefore considerably longer than the critical wavelength of a hollow tube of the same outer dimensions.

Flexible wave-guides

23. A great disadvantage in the use of ordinary wave-guides is their non-flexibility and types of flexible circular wave-guide have been designed. They consist of two helical strips of metal, usually copper-bronze or galvanized steel, which are interlocked to form a complete cylindrical pipe. The surface of the wave-guide is ribbed and has the appearance of "vacuum cleaner" hosing. In order to obtain good electrical contact between the strips, the wave-guide is tightly constructed and this impairs the flexibility of the wave-guide to a great extent.

WAVE-GUIDE ATTENUATION

24. As stated in para. 20 the main reason for the use of wave-guides is because the attenuation is negligible compared with that of feeders when used at very high frequencies. For distances of the order of 3 ft. the attenuation of a feeder does not seriously affect the

amplitude of a signal passing along it, and the convenient size and flexibility make it more advantageous to use than a wave-guide. On the other hand, it is sometimes necessary to transmit a signal over a distance of the order of 20 ft. or more, in which case the attenuation of a feeder would be prohibitive to its use; a wave-guide using air as dielectric, however, has negligible attenuation over lengths of the order of 20 ft. and this property far out-weighs the inconvenience due to size and non-flexibility when in use over such distances.

25. Concentric feeder may be of the flexible type, the inner conductor being embedded in a high quality dielectric such as distrene and the outer being in the form of metal braiding; or it may be rigid, the inner being supported by means of insulating discs. Commercial dielectrics are found to have very heavy losses at hyper-frequencies, these losses increasing with increasing frequency. Thus, flexible concentric cable can only be used at short lengths. Rigid feeder has the minimum of supporting dielectric (each support is honeycombed as much as the necessary supporting strength will permit) so that there is a minimum of dielectric loss; each of the supports, however, produces a mis-match in the line, so that some of the signal will be reflected back. This can be overcome by careful spacing of the supports, but the feeder attenuation then becomes extremely sensitive to frequency.

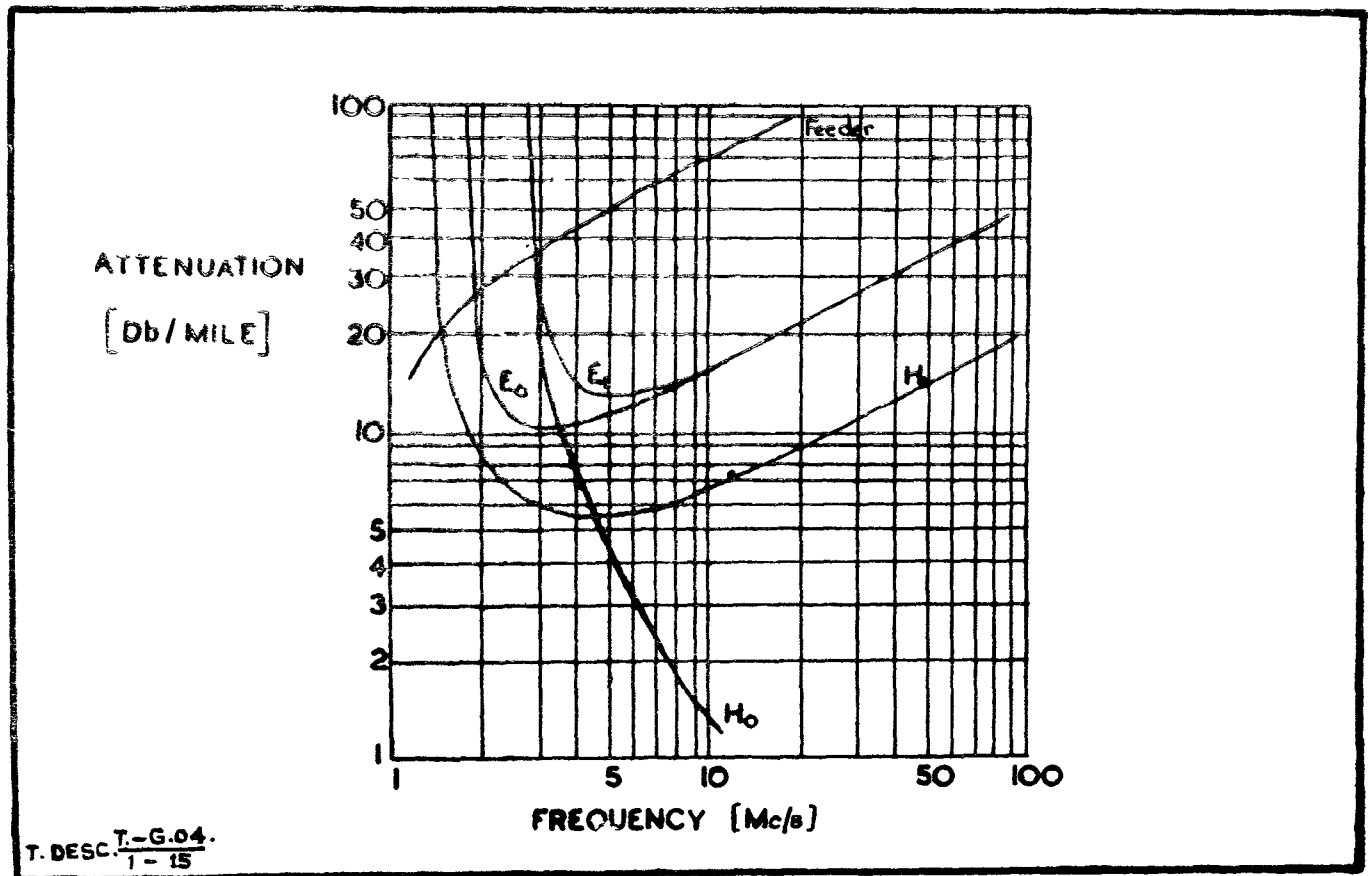


Fig. 15

26. Mathematical calculations have been made on the attenuations of a wave-guide and concentric feeder due to the ohmic losses alone in the conductors themselves and the results have been verified practically. Figure 15 indicates attenuation for various modes in a circular wave-guide and also for concentric feeder, plotted against frequency on a logarithmic scale.

27. It can be seen that, whereas the attenuation in a wave-guide generally falls to a minimum from a high value at the cut-off frequency, the feeder attenuation rises steadily with frequency, being somewhat larger than that for a wave-guide. Figure 15 also shows that H_0 waves possess the remarkable property that the attenuation does not fall to a minimum but decreases indefinitely with increasing frequency.

28. When using wave-guides, care must be taken to see that the inner surface does not contain irregularities, since they will produce distortions in the wave-pattern, and also give back-reflections, so that the attenuation will be increased. If the inside is damp, the currents induced in the wave-guide wall will reside mainly in the coating film of water; this will again increase attenuation a great deal owing to the poor conductivity of the film.

29. The attenuation in flexible wave-guides is somewhat larger than in rigid ones since the ribbed construction of the inner surface produces back-reflections and also the contact between the helical strips is not perfect, with the result that ohmic losses are increased.

WAVE-GUIDE IMPEDANCE

30. Experiments on wave-guides seem to indicate that they have many of the properties of transmission lines. For example, a resistive film placed over the mouth of a wave-guide and backed by a short circuited quarter-wavelength of guide will absorb electromagnetic energy which is passing down the wave-guide; generally, there will also be a back-reflection of energy along the guide, producing standing waves. If, however, the resistance per centimetre square has a particular value, then all the energy will be absorbed and the film will "match" the wave-guides; this particular value corresponds to the characteristic resistance in the case of a transmission line.

31. Characteristic impedance of a transmission line is an extremely important and useful quantity and it might be expected that a similar quantity is equally important in wave-guide theory. For transmission lines, it was found much easier to develop the theory from observations on currents and potentials along the line, so that characteristic impedance was defined as the ratio of potential difference to current at any point in an infinite line. In the case of wave-guides, however, the currents and voltages are distributed

and it was found much easier to develop the theory from observations on the electromagnetic fields in the wave-guide; a new quantity known as the *characteristic wave-impedance* was therefore defined to correspond with characteristic impedance in the current-potential theory.

32. Consider a three-dimensional orthogonal co-ordinate system OX, OY, OZ and let the axis of an infinite wave-guide of any cross-section be parallel to OZ as shown in figure 16.

Then the characteristic wave-impedance Z is defined as

$$Z = \frac{E_x}{H_y} = -\frac{E_y}{H_x} \dots \dots \dots (4)$$

$$\text{For E waves, } Z = \sqrt{\frac{\mu}{\epsilon}} \cdot \frac{\lambda_g}{\lambda_a} \dots \dots \dots (5)$$

$$\text{H waves } Z = \sqrt{\frac{\mu}{\epsilon}} \cdot \frac{\lambda_g}{\lambda_a} \dots \dots \dots (6)$$

33. The practical units of wave-impedance are found to be "ohms per centimetre square" and it is found that if the wave-guide is terminated by a load which has a resistance per unit area equal to the characteristic wave-impedance, then a perfect match is obtained. This shows the close analogy between wave-guides and feeders.

34. As the dimensions of the wave-guide are increased, λ_g approaches λ_a and Z as given by equations (5) and (6) approaches the value $\sqrt{\mu/\epsilon}$. This is known as the "Intrinsic Impedance" of the medium and is the ratio of E to H in an infinite plane wave in the medium. The intrinsic impedance of free space itself is calculated to be 377 ohms.

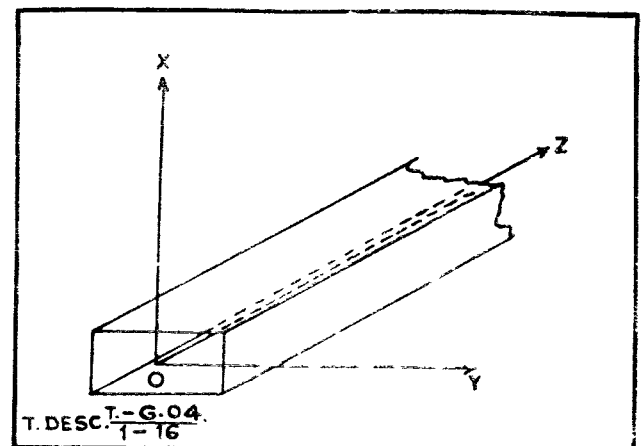


Fig. 16

(51553/127/5/MGME-12)

END

THE MAGNETRON VALVE

General Principles

[Based on E. M. E. R. Telecommunications A 012, Issue 1.]

Limitations of normal valve circuits.

1. Normal valve circuits are not satisfactory for use at wave-lengths under about one meter and are quite impracticable in the centimetre range (λ of the order of 10 cms.). Not only is the construction, from lumped impedance, of a circuit to resonate at 3,000 Mc. extremely difficult, but valve constants render normal valves incapable either of generating or of amplifying such frequencies.

2. The low performance of normal valves is due to three main causes, viz. :—

- (a) valve lead inductances.
- (b) inter-electrode capacitances.
- (c) transit-time effects.

3. The combination of (a) and (b) might produce an oscillatory circuit which would entirely mask the effect of any external tuned circuit. Cathode lead inductance produces an effective grid input resistance which is inversely proportional to the frequency squared. Anode and screen lead inductances are found to give rise to instability.

4. (a) The most serious consideration is, however, the transit-time effect. At frequencies of the order of 3,000 Mc. the transit-time, i.e., the time taken for the electron to travel from the cathode to the anode under the influence of the anode voltage, can be shown to be of the same order of magnitude as the periodic time of the oscillations. This leads to a very low grid input resistance, again inversely proportional to the square of the frequency. Table I shows approximate values of grid input resistance for a Mullard EF 50* at various total valve currents.

Frequency	Grid input resistance.		
	Valve current 2Ma	Valve current 5Ma	Valve current 10Ma
20 Mc.	125k Ω	100k Ω	60k Ω
50 Mc.	25k Ω	14k Ω	10k Ω
100 Mc.	6.3k Ω	3k Ω	2.5k Ω
200 Mc.	1.5k Ω	0.8k Ω	0.65k Ω
500 Mc.	0.25k Ω	0.16k Ω	0.10k Ω

Table 1. Grid input resistance, at various frequencies, of a Mullard EF 50 valve.

The grid input resistance becomes so low at ultra-high frequencies that the damping on the previous stage reduces its gain to less than unity. At these frequencies, normal valves cannot, therefore, be used as amplifiers.

(b) In addition to the low grid input resistance, there is a phase difference between the input to the grid and the anode output, due to the finite transit-time from grid to anode. This phase difference is dependent upon the anode voltage. The difficulty of feeding energy back into the grid circuit from the anode circuit, together with the low grid input resistance, makes the normal valve highly unsatisfactory as an oscillator.

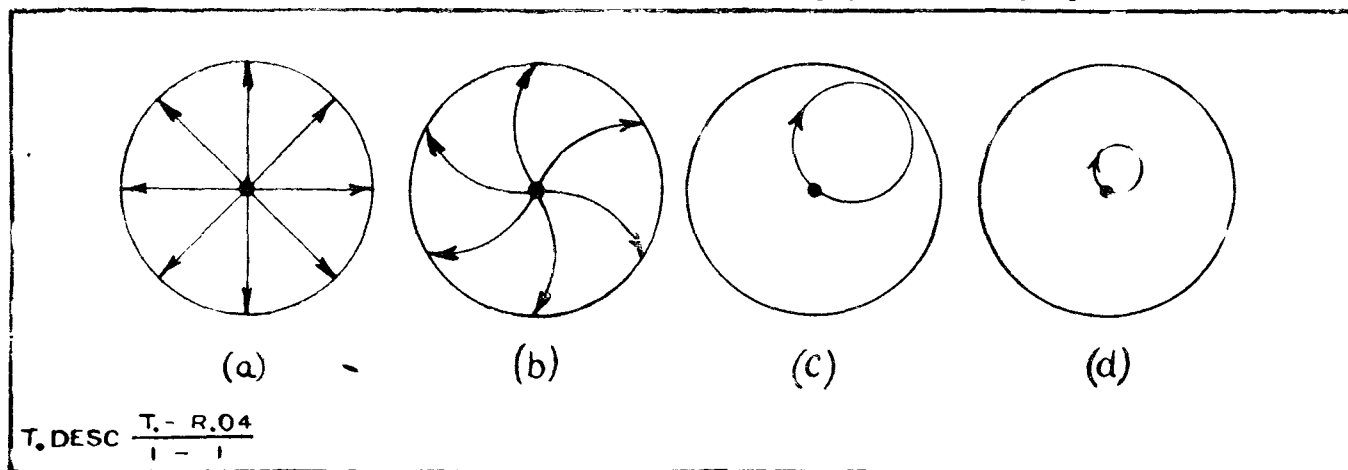


Fig. 1. Electron paths for various magnetic field strengths.

5. The original types of oscillator used for the generation of U.H.F. waves were the Barkhausen-Kurz positive-grid oscillator and the magnetron. The latter has been developed in recent years and forms one of the fundamental valves used in centimetric technique.

The diode in an axial magnetic field

6. As a preliminary to the study of the magnetron consider a cylindrical diode with a thin cathode at the centre and a magnetic field along the axis. Let the anode potential be V , the magnetic field strength H , and the radius of the anode a .

7. Under the combined influence of the magnetic and electric fields, the path of the electron will be curved. If the space charge be neglected, it can be shown that the path is circular and, for a given voltage, that the radius will be inversely proportional to H .

Fig. 1 shows the electron paths for various strengths of the magnetic field. In fig. 1 (a) the magnetic field strength is zero and the electron paths are radial. Fig. 1 (b) shows the bending due to a small magnetic field; here the diameter of the circular electron path is greater than the anode radius. The effect of very high magnetic field is shown in fig. 1 (d). There is obviously a critical value of the magnetic field at which the electron just misses the anode. Variations of the magnetic field near this critical value will produce large variations of anode current; if the critical value is exceeded, the current will drop to zero [fig. 1 (c)].

8. The value of the critical or cut-off field H_c depends on the anode radius and voltage and is given by the equation,

$$H_c = \frac{6.72\sqrt{V}}{a} \text{ gauss}$$

where V = anode volts

a = anode radius in cm., assuming that it is large compared with the filament radius.

9. (a) The above arguments presuppose the absence of space charge near the cathode. In practice, this condition exists only when the valve is first switched on.

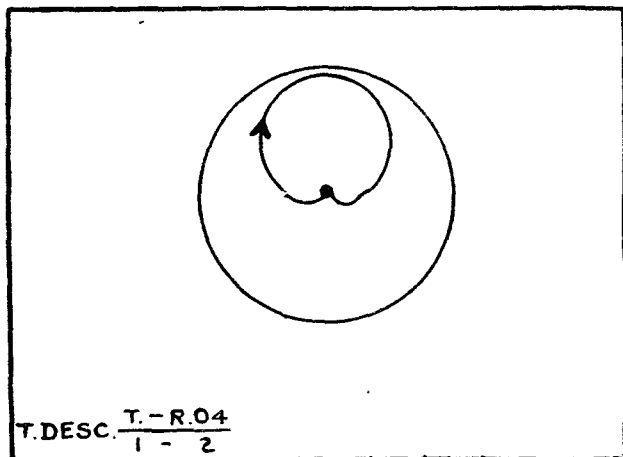


Fig. 2. Electron path neglecting effect of magnetic field on space-charge distribution.

*Figs. 4 and 5 from "Thermionic Tubes at Very High Frequencies" by A. F. Harvey. Chapman & Hall, 1941

Under normal conditions of operation, space charge will be present and this modifies the circular path as in fig. 2.

(b) If the effect of the magnetic field on the space charge is taken into account, the picture is still further modified. Fig. 3 shows a comparison of the two cases. Curve 1 is the path neglecting the effect of the magnetic field upon the space charge distribution, while curve 2 takes it into account.

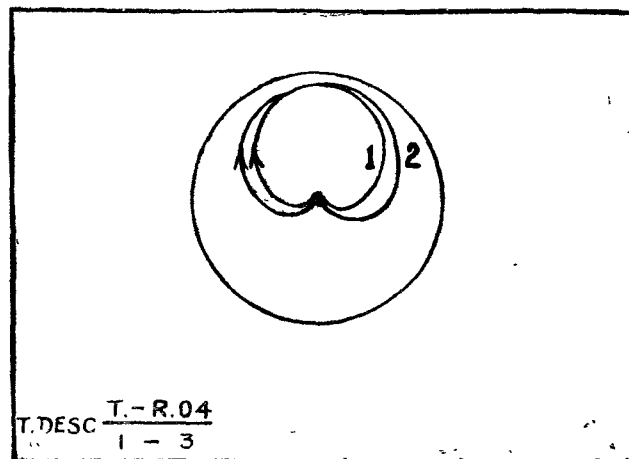


Fig. 3. Electron paths compared, with and without consideration of effect of magnetic field.

10. Figs. 4 and 5* respectively show curves of plate current against magnetic field (voltage constant) and plate current against anode voltage (magnetic field constant).

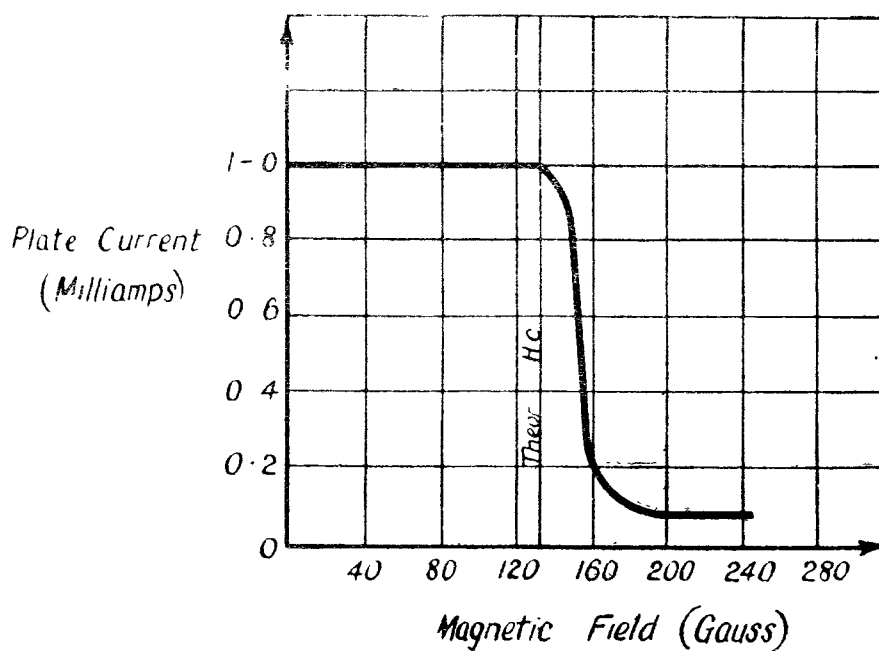
11. It will be noticed that there is a discrepancy between the cut-off field calculated and that which is found and that there is a considerable rounding both at the onset and the end of the drop in anode current. There are various causes for this:—

- Inaccuracy in lining up the magnetic field along the axis of symmetry of the diode. This is highly important and the effect of tilt of field is shown in fig. 6.
- Eccentricity of the cathode has a large effect and is obviously linked up with (a).
- The presence of high frequency oscillations of small amplitude has considerable effect.
- The different emission velocities of the electrons from the cathode have a small effect.
- A second order effect is due to the voltage drop along the filament.

The split-anode magnetron as an oscillator.

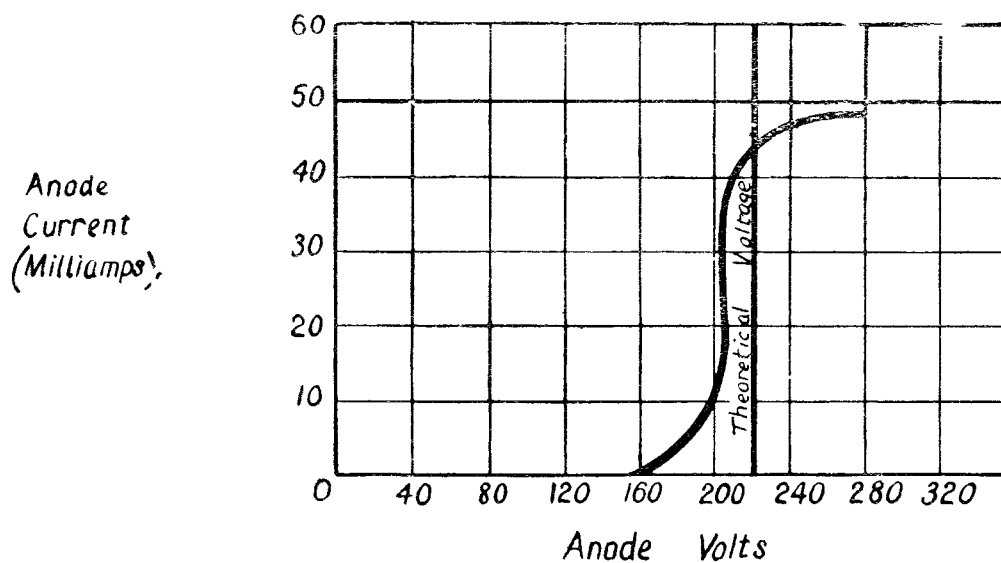
12. It has been found that a cylindrical diode having its anode split into two or more segments can function as an oscillator under the influence of an axial magnetic field. The modes of oscillation fall into three main divisions, which can be summarised as follows:—

(a) *The dynatron mode.*—The magnetic field required is equal to, or greater than, cut-off. The wave-length is dependent on the external tuned circuit and is dependent on the dimensions of the tube and of the electric and



T. DESC. $\frac{T - R 04}{1 - 4}$

Fig. 4. Curves of plate current against magnetic field strength, with plate voltage constant.



T. DESC. $\frac{T - R.04}{1 - 5}$

Fig. 5. Curves of plate current against plate voltage, with magnetic field strength constant.

DESCRIPTION
TELS. R. 04

magnetic fields. The only frequency limit is the upper one; it is necessary that the periodic time be much greater than the electron transit-time.

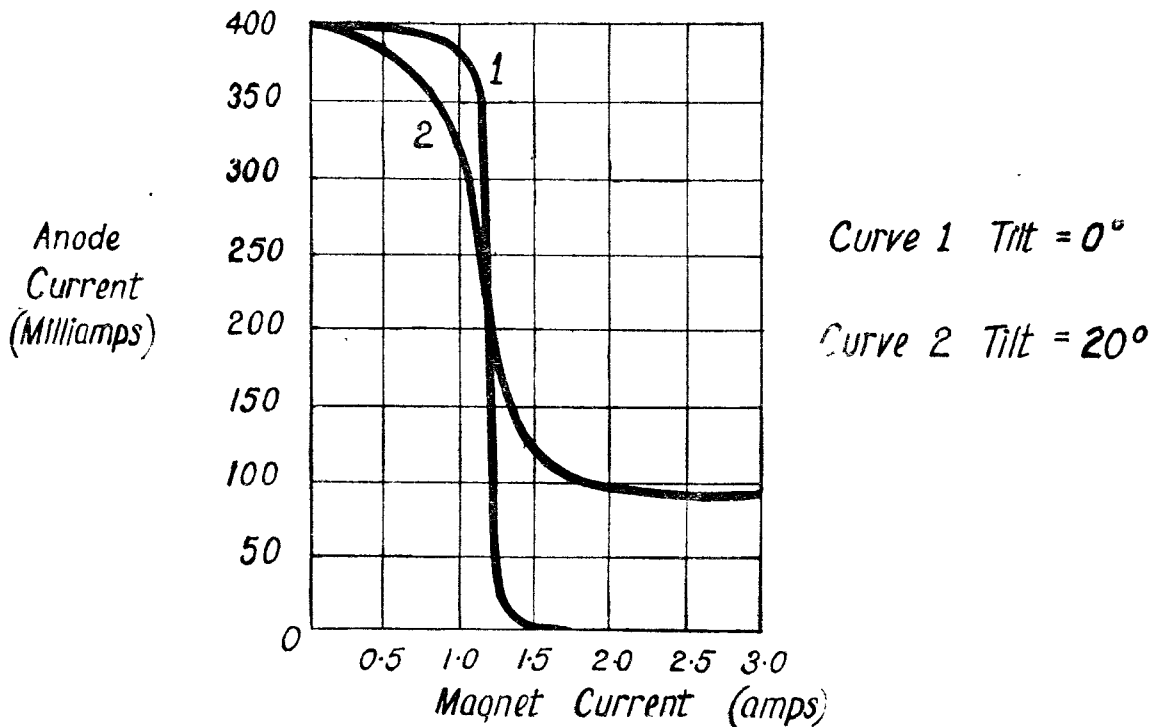
(b) *The resonance mode.*—The magnetic field required is greater than cut-off. The wave-length is substantially independent of the external circuit and is a function of the tube dimensions and field strength. The

D. M. E. (INDIA)

TECHNICAL INSTRUCTIONS

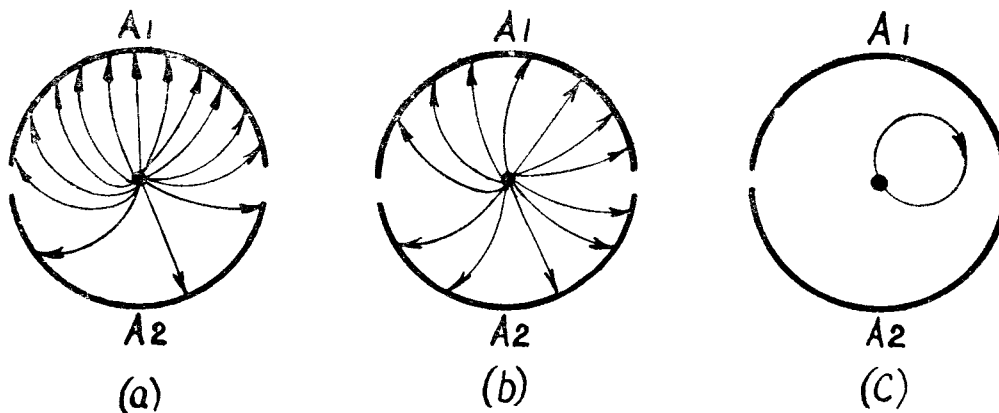
wave-length is proportional to the magnetic field for a given tube and anode voltage. Wave-lengths range from 10 cm. to indefinitely long wave-lengths.

(c) *The electron mode.*—The magnetic field must be near cut-off. The wave-length is independent of external circuits and is usually less than 15 cm. Wave-length is inversely proportional to magnetic field.



T.DESC. T.-R.04
1-6

Fig. 6. Effect of tilt of magnetic field on characteristic.



T.DESC. T.-R.04
1-7

Fig. 7. Electron paths in split-anode magnetron.

The dynatron mode.

13. (a) Consider a cylindrical diode in the presence of an axial magnetic field as before, but let the anode be split into two parts A_1 and A_2 (fig. 7). Let A_1 be at voltage V and let A_2 be at zero potential, while the magnetic field is such as to cut off the anode current if both segments were at V . The electrostatic field opposes the deflecting effect of the magnetic field and causes a large current to flow to A_1 but only a small (almost zero) current to flow to A_2 (fig. 7a). Now, as the voltage on A_2 is increased, the electrostatic field becomes less operative and more current flows to A_2 at the expense of the current in A_1 (fig. 7b). When the voltage on A_2 approaches that on A_1 the current in both anodes starts to fall rapidly, since the system is approaching the cut-off condition. The current on both anodes will be zero when A_2 is at voltage V (fig. 7c). Further increases in the voltage of A_2 will serve to increase the current to both anodes, since the magnetic field becomes insufficient to maintain the cut-off conditions. It is, however, obvious that A_2 will now draw the larger current, since it is at the higher potential. Fig. 8 gives a schematic representation of the currents to be expected.

(b) There is a region of negative resistance where increase in anode voltage causes decrease in current near the critical value of voltage. In fig. 9 curves are given showing how the currents in the two segments vary as either anode is varied about a mean for two different values of magnetic field.

It can be seen that, within limits, the current flowing to the segment at the lower potential is greater than the current flowing to the segment at the higher potential.

The range of this negative resistance region is increased by using a higher value of magnetic field.

(c) Any tube which has negative resistance properties is capable of functioning as an oscillator when connected to a tuned circuit of sufficient dynamic resistance. The circuit of fig. 10 will then function as an oscillator, the peak-to-peak voltage being such that the anode voltage swings within the negative resistive range.

(d) With a high value of magnetic field, current will flow only on the peaks of the oscillation and the magnetron acts as a class C push-pull dynatron oscillator. The ordinary treatment of a dynatron oscillator can be applied to this type of circuit.

(e) The efficiency of this long-wave dynatron mode of the magnetron is rather less than that of a tetrode of the same anode voltage and filament emission; it falls almost to zero when the electron transit-time becomes comparable with the periodic time of the oscillation to be maintained.

(f) It is to be noted that the values of the electric and magnetic fields required do not depend on the frequency to be generated; thus, provided that the magnetic field is equal to or greater than cut-off, any frequency can be maintained, subject to the higher limit already quoted, without alteration of field strengths.

The resonance mode.

14. (a) For values of magnetic field greater than cut-off, a type of oscillation occurs in which the wave-length is not very dependent upon the constants of the external circuit. This is quite clearly due to causes different from those just described in the dynatron mode. The wave-length is actually dependent on the field.

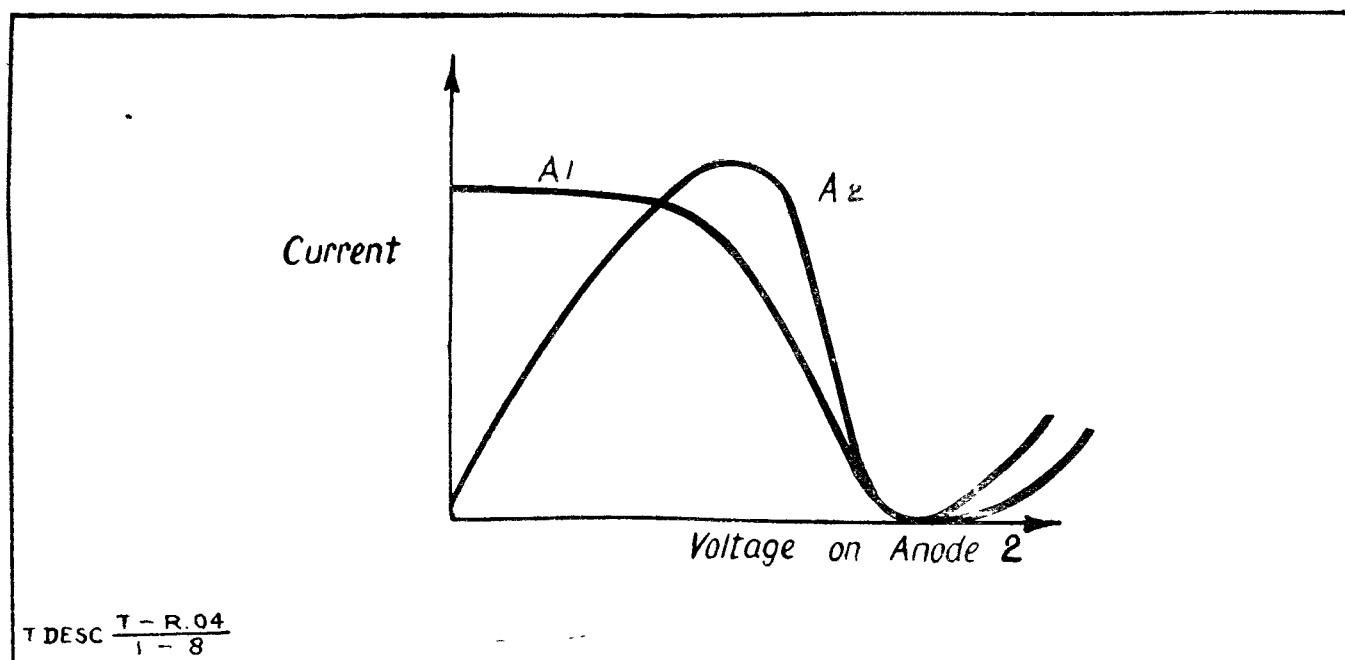


Fig. 8. Curves of plate currents against plate voltages.

DESCRIPTION
TELS. R. 04

D. M. E. (INDIA)

TECHNICAL INSTRUCTIONS

strength, and for any given set of conditions :—

$$\omega = \frac{2V_k 10^8}{a^2 H}$$

where $\omega = 2\pi f$ (f =frequency).

V = mean anode voltage.

H = magnetic field.

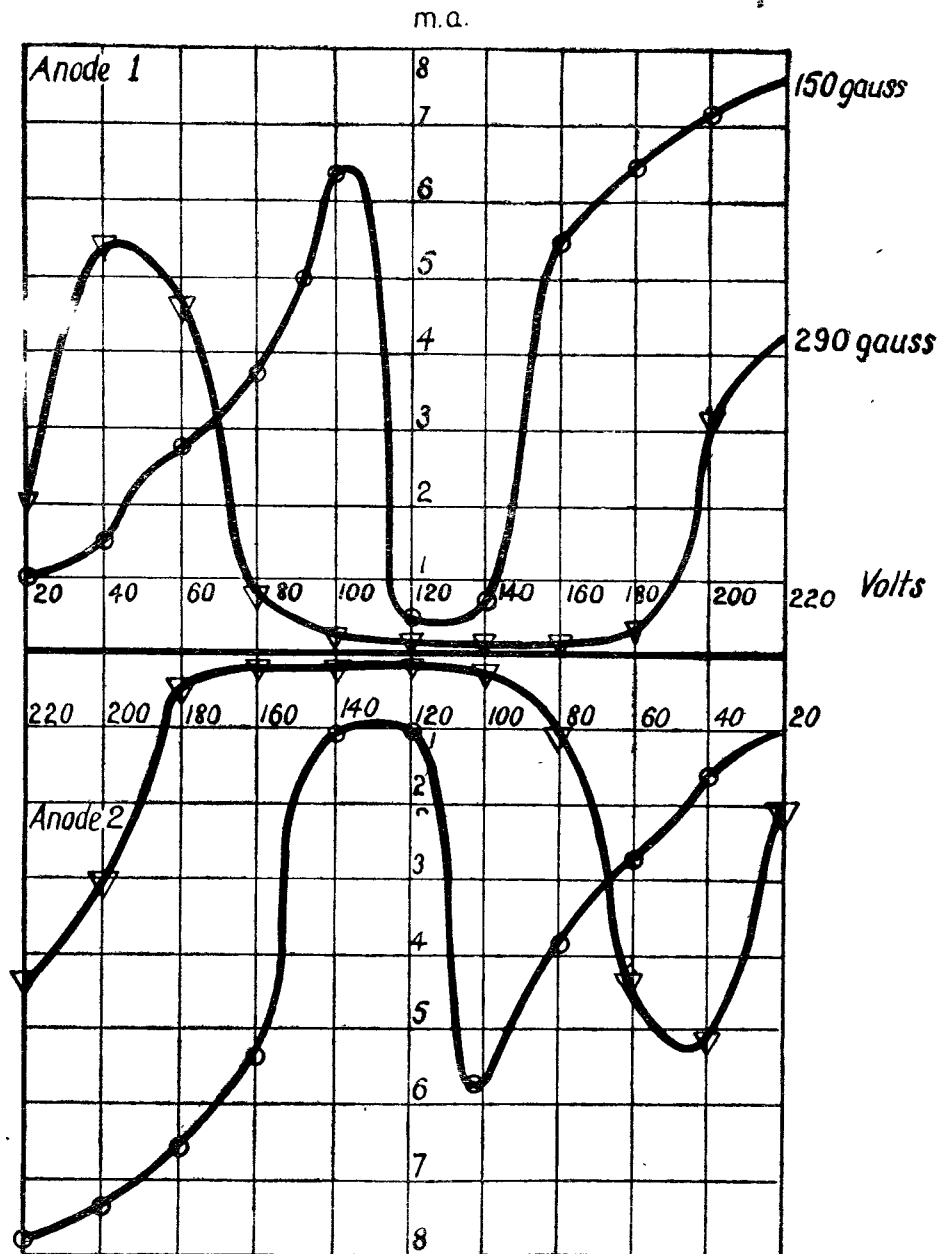
a = anode radius.

k = number of pairs of segments.

The formula gives the optimum wave-length for the conditions and only a small variation can be made, by detuning the external circuit, without causing the oscillations to cease. For any given magnetron :—

$$\frac{\lambda V}{H} = \text{constant}$$

where λ = wave-length.



T.DESC. $\frac{T.-R.04}{1-9}$

Fig. 9. Curves of plate currents against variations in plate voltages and magnetic field strength.

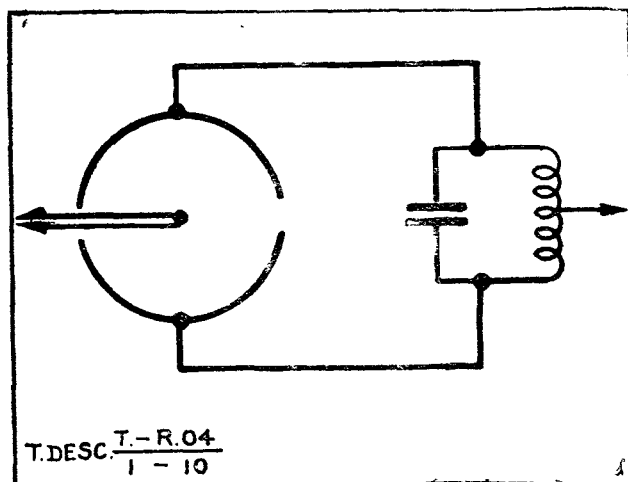


Fig. 10. Split-anode magnetron connected as an oscillator.

(b) Wave-lengths from the order of 10 cm. up to indefinitely long wave-lengths are obtainable.

(c) In a cylindrical magnetron, most of the voltage drop occurs near the filament and only a small residue spreads across the space to the anodes. The motion of the electrons in the field has been the subject of much study. It is suggested that if the anodes are subjected to an alternating voltage $V = v \sin \omega t$, then an electron passing one of the gaps will suffer deflection towards the anode at the lower potential. If this has the effect of causing the electron to hit that anode, the dynatron oscillation occurs. If, however, the deflection is insufficient for this, the electron will carry on along its path towards the cathode; but, instead of reaching the cathode, the deflection suffered will cause it to miss the cathode and a type of curved cycloidal path will be followed. Provided the alternating anode voltage is of the correct frequency, then at each passage past a gap the electron will suffer a further deflection and deceleration until it ultimately hits one of the anodes (fig. 11).

(d) Support for this theory is given by G. R. Kilgore who, by using a magnetron with a small amount of

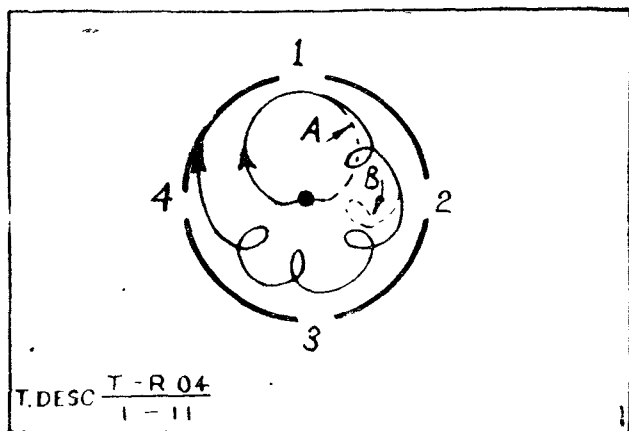


Fig. 11. Cycloidal electron path with alternating voltage applied to anodes.

residual gas present, was able to photograph the ionisation track left in the path of the electron.

The dotted lines in fig. 11 show the path the electron would have taken had it not been deflected by gap 1 (A) and subsequently by gap 2 (B).

(e) Measurements made on a magnetron when oscillating in this mode showed that a resonant condition existed in the tube, which had the usual features of a series circuit, but with component impedances consisting of negative resistance, negative inductance and negative capacitance, giving rise to the name "resonance mode".

(f) From the formula $w = \frac{2V_k 10^8}{a^2 H}$ it is seen that the wave-length is proportional to the square of the anode radius and inversely proportional to the number of segments. Thus, for ultra-short wave-lengths, a small diameter anode cut into a large number of segments is needed. The consequent limiting of the permissible dissipation makes the tube capable of handling only small powers.

The electronic mode.

15. (a) This mode is characterised by the generation of a wave-length that is closely related to the electron transit-time and is practically independent of external circuit constants. Wave-lengths obtainable in this mode are usually less than 15 cm. and occur in the neighbourhood of the critical magnetic field. Experiment has shown that wave-length is approximately given by the formula: $\lambda H = 11,000$. This is in close agreement with the assumption that the oscillations have a periodic time equal to twice the calculated transit-time of an electron from cathode to anode.

(b) In general, it is found that the angle between the magnetic field and the electrode axis is critical and is not zero. The width of the anode gaps has no important effect on efficiency, which is greatest with small anode currents and falls rapidly when the space charge density becomes large. This suggests that it is not the arrival of electrons at the anode, but the transfer of energy from the electrons to the circuit during transit, that determines whether or not oscillations can be maintained.

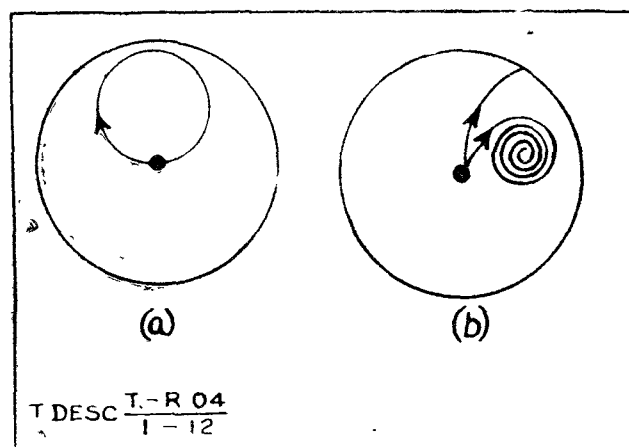


Fig. 12. Electron paths with critical anode voltages.

DESCRIPTION

TELS. R. 04

(c) From the above facts it is possible to form an explanation of how the oscillations are maintained. Since the width of the gaps has no effect, the case of a cylindrical diode having a voltage just less than the critical value can be considered, so that, for a steady anode voltage, the electron path will be a circle (fig. 12A). It should be understood that, in practice, the anode will always be split into two or more segments to enable a tuned circuit to be incorporated.

(d) Now consider that an alternating potential of the electronic frequency is superimposed on the D.C. anode potential. An electron which leaves the cathode during the first part of the positive half-cycle will receive a little extra energy and will hit the anode, where its energy is dissipated as heat. An electron leaving the cathode about a half-cycle later, when the anode is going negative, will give up energy to the anode circuit, for, whether it be approaching the anode or receding from it, the anode alternations are always such as to oppose the radial motion of the electron. The electron path becomes a closing spiral. As the spiral gets smaller the electron circle takes place more and more in the region of almost constant potential near the anode and its rotation frequency increases. Thus, the phase of the useful group of electrons relative to the alternating anode voltage changes continuously until, instead of giving up energy to the alternating circuit, the electrons start to abstract energy, with a consequent opening of the spiral, until they reach the anode with almost their original velocity. The net output will therefore be very small.

(e) The effect of tilting the magnetic field is to pull the spiral out to a helical path, of which the axis tends towards the anode, along the magnetic field. If the tilt is critically set, the electrons can be made to hit the anode just as their period of usefulness is over, thus giving optimum efficiency. The tilt of the field has also the effect of eliminating undesired long-wave oscillations of the other modes. The optimum tilt

D. M. E. (INDIA)

TECHNICAL INSTRUCTIONS

depends on the filament emission as shown in fig. 13 for a typical magnetron.

(f) The magnetron in the electronic mode is capable of producing extremely short wave-lengths, although power output and efficiency are low. The optimum anode voltage and filament emission increases with frequency; thus, the wave-length is limited by the allowable dissipation.

Frequency drift in the electronic mode.

16. (a) Frequency drift in the magnetron is due to:—

(i) temperature changes as the valve warms up after switching on.

(ii) variation of anode and filament voltages consequent on mains supply fluctuations.

(b) The greater part of the drift due to temperature changes occurs during the first two minutes, after which the frequency is within 2 Mc. of its final value. After a further ten minutes the frequency is within 1 Mc. of its final value.

(c) Frequency change with filament voltage is negligible within quite wide limits of variation of the heating voltage (fig. 18).

(d) The drift with anode voltage is considerable, although the drift decreases as the loading of the H.F. circuit increases. It is, however, possible to choose a voltage (fig. 17) such that a small fluctuation of mains voltage causes only a small frequency variation. In general, a fairly well stabilised supply is satisfactory.

Conclusion.

17. It has been assumed throughout that an external tuned circuit is attached to the anodes, but this is not always so. It is possible to choose the dimensions of the anode, as regards length, radius, and distance apart of the segments, so that the whole anode system is the resonant circuit and no external circuit is then necessary. Each pair of anode fingers can be considered to

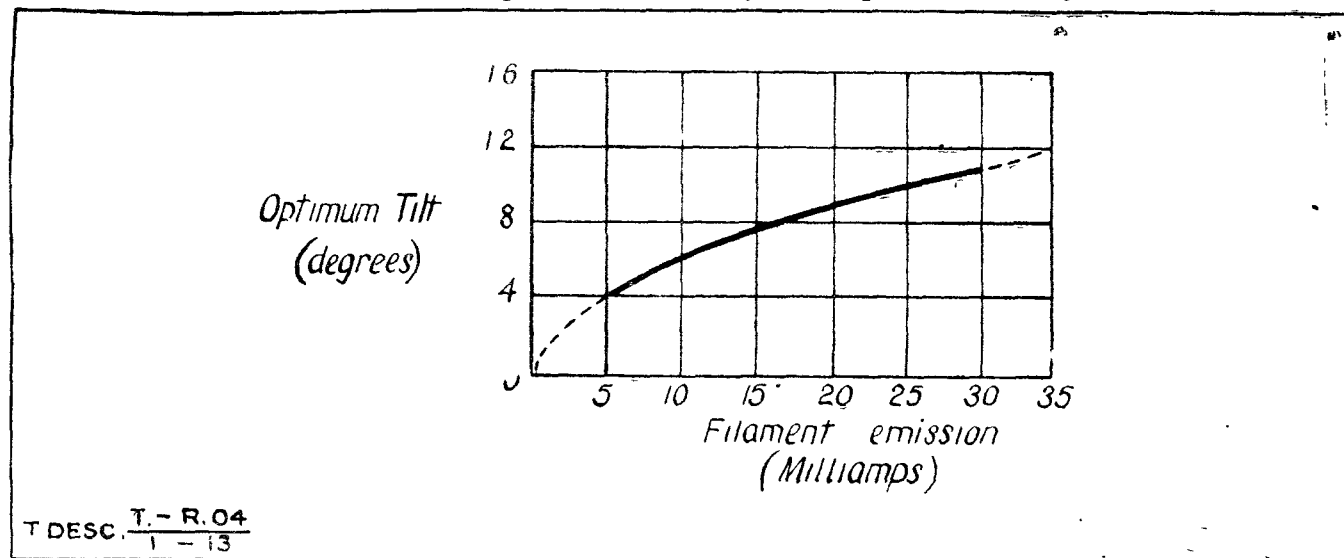


Fig. 13. Effect of tilt of field on emission.

TECHNICAL INSTRUCTIONS

form a section of transmission line, short-circuited at one end, behaving as quarter-wave resonators. This gives rise to the term "resonant segment magnetron". Fig. 14 shows a diagrammatic sketch of the G.E.C. magnetron, type E. 1210B, which is fitted with a standard octal base. Figs. 15 to 18 show the effect on wavelength of the variables. The curves refer to the electronic

mode. The discontinuity in fig. 15 is due to a change of mode of oscillation of the lechers.

18. As yet there is no complete theory to explain all the phenomena associated with the magnetron, but there is some evidence that the modes of oscillation are not so sharply defined as this treatment suggests.

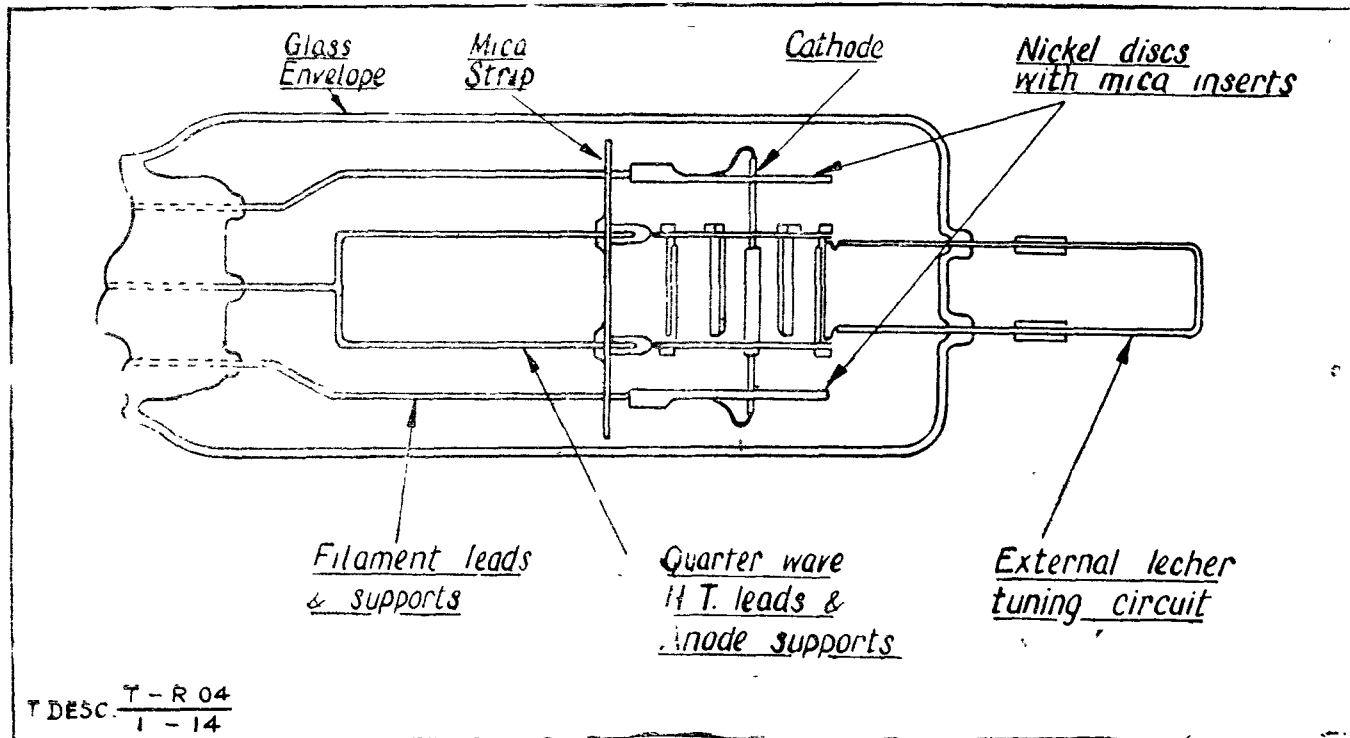


Fig. 14. Structure of G.E.C. magnetron type E.1210B.

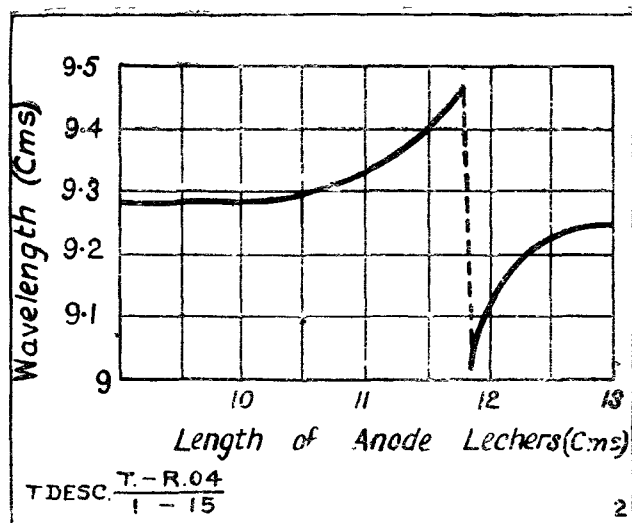


Fig. 15.

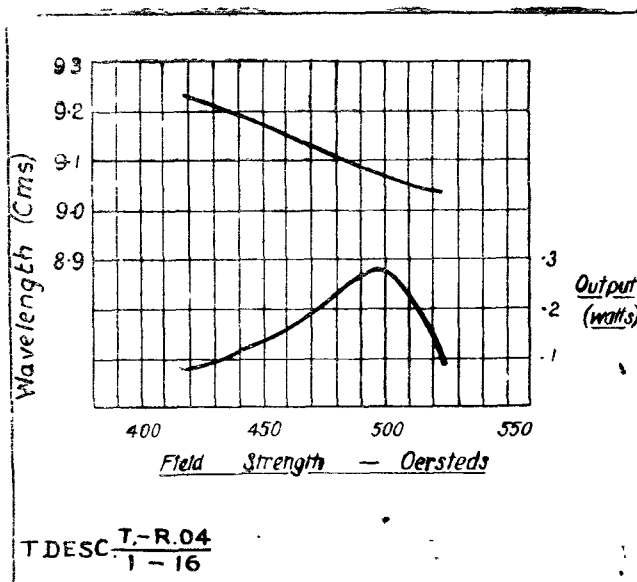


Fig. 16.

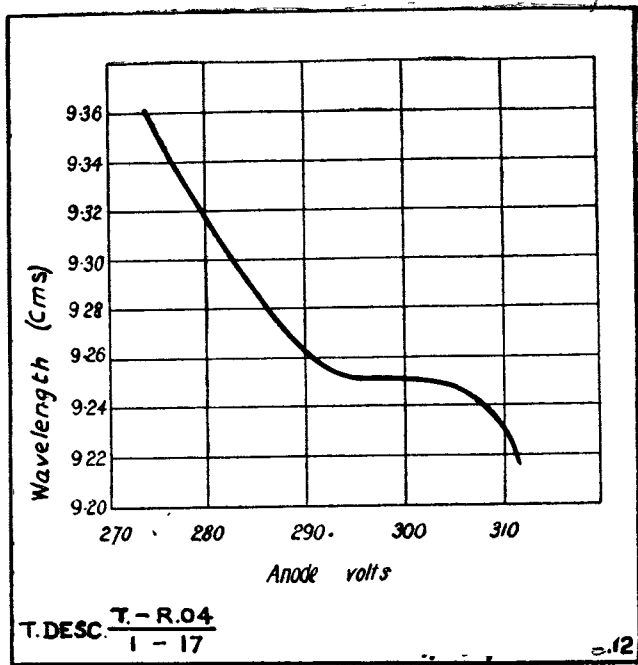


Fig. 17.

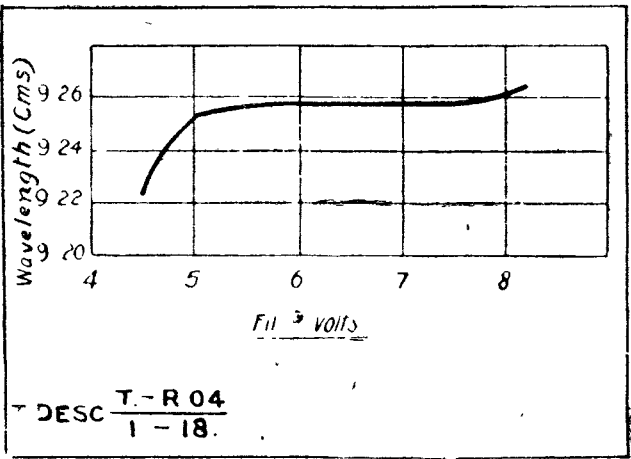


Fig. 18.

Figs. 15—18. Effects of variable factors on wave-length generated.

END

(51549/24/MG-ME12.)

TRAINER RADAR A. A. NO. 1

General Description

[Based on E. M. E. R. Tels. S. 542, Issue 1, date]

INTRODUCTION

Function

1. The GL 4A trainer is designed to operate in conjunction with A.A. No. 1 Mk. II equipments operating in the 55–85 Mc/s band. By the use of this equipment an echo can be produced on the CR tubes of the A.A. No. 1 receiver similar to that from an aircraft in flight. Continuous change of apparent range between the limits of 5,000 and 30,000 yards may be effected, but the elevation remains constant at about zero.

Brief description

2. The equipment consists essentially of :—

- (a) A square-topped wave generator (a transitron), triggered by a negative locking pulse.
- (b) A differentiating circuit of short time constant to differentiate the square-topped wave into a negative and a positive pulse.
- (c) A Westector to discriminate against the negative pulse.
- (d) An output oscillator covering the frequency range 55–85 Mc/s.

3. The output oscillator is normally biased to cut-off and oscillates only when driven by the positive pulse from the transitron. A pulse is thus radiated a short time after the incidence of the negative lock pulse ; the actual delay time is equal to the width of the square-topped wave from the transitron and is variable in time between predetermined limits.

4. The negative locking pulse can either be fed in *via* a three-way cable from the A.A. No. 1 Mk. II transmitter, or be produced internally by means of a gas discharge triode. The whole equipment is housed in a metal box, which, together with the aerial, can be carried on a man's back.

TECHNICAL DESCRIPTION

Electrical circuit

5. The circuit diagram is shown in fig. 1. V2 is the transitron and is provided with anode load R5 and screen load R2A. The screen is coupled to the suppressor by condensers C5, which is fixed and C11, which is variable and in parallel with C5. C11 is driven by the motor M *via* a step-down gear box at a speed of about 1 revolution in 5 minutes. Resistors R12 and R9B are included in the leads to the condenser C11 to prevent excessive feed lock when C11 is approaching maximum capacity.

6. The suppressor is returned to earth *via* R9 and R8 the latter being variable. The control grid is biased by the voltage drop across R6, and the suppression by the voltage drop across R6 and R7.

7. Under these conditions, the valve V2 (VR65) is normally conducting, until the negative pulse, which is derived either from the transmitter or from the internal pulse generator, is applied to the grid *via* C4. This

causes an immediate fall in the cathode current and consequently a rise in the potentials of the anode and the screen, which causes the suppressor to acquire a positive potential. Electrons collect on the suppressor and at the end of the grid pulse, the suppressor is left more negative than originally. This has the effect of reducing the anode current and increasing the screen current. The resulting drop in screen potential is passed to the suppressor *via* C5 and C11, the effect being cumulative until the anode current is cut off.

8. The negative charge on the suppressor now leaks away through R8 and R9 until anode current flows again, when the valve abruptly returns to its former condition. A positive square-topped wave is thus produced at the anode and a negative square-topped wave at the screen, the duration of the wave being determined by the time constant of C5, C11, R8 and R9 and being varied by either C11 or R8.

9. The square-topped wave at the screen is now differentiated by C6 and R9A (this combination having a short time constant) into a negative and a positive pulse. The onset of the negative pulse coincides with the onset of the transitron pulse and the onset of the positive pulse with the termination of the transitron pulse, the time delay between the negative and positive pulses being equal to the width of the square wave.

10. Across the resistance R9A is shunted a Westector which discriminates against the negative pulse but allows the positive pulse to pass through C6A and L2 to the grid of V3 which is arranged as a choke-fed Colpitts oscillator, and is normally biased to cut-off. The positive pulse then allows V3 to oscillate until it biases itself off by the development of grid current. Before this negative charge can leak away and allow the valve to oscillate again, the positive pulse terminates and the valve is left in the current cut-off condition until the incidence of the next pulse. The oscillator V3 is arranged to cover the band 55–85 Mc/s, the frequency being mainly determined by L3 and C8.

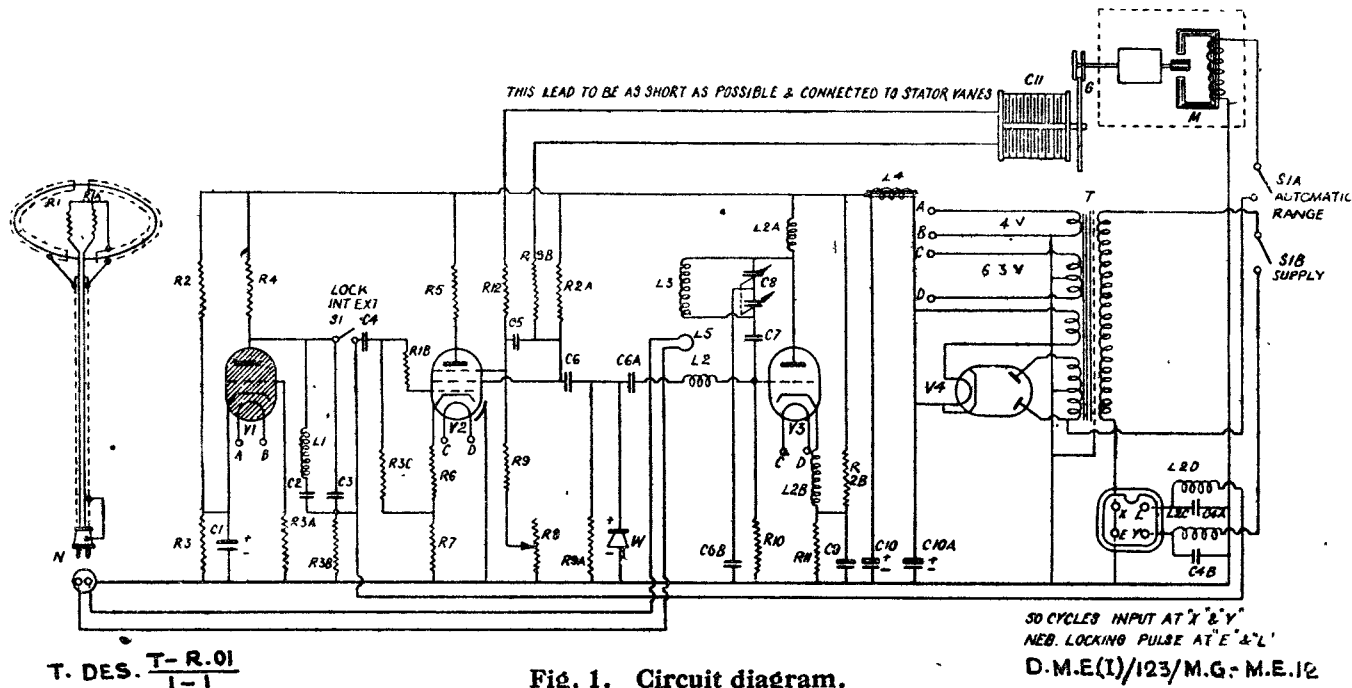
11. A single-turn coupling coil L5 is spaced about $\frac{1}{2}$ in. from the tuning coil L3 and couples the oscillator *via* the twin wire feeder to a semi-aperiodic aerial, which covers the required frequency band without tuning adjustments. The radiation from this aerial is therefore in the form of pulses at a recurrence frequency equal to that of the A.A. No. 1 Mk. II transmitter or to that of the internal pulse generator (whichever is in use) and delayed to an extent dependent upon the instantaneous time constant of C5, C11, R8 and R9.

12. With the switch S1 in the up or open position, the trainer will operate without the main transmitter as a source of locking. The gas triode V1 produces a negative pulse, at the junction of R3B, C2 and C3, which can be used both for triggering the transitron V2 and for synchronizing the time bases of the receiver display circuits.

13. V1 operates as a saw-tooth oscillator, the valve being normally cut off and only conducting upon ionization, which takes place when the anode voltage reaches a positive value about twenty times that of the negative grid bias. When the valve ionizes and strikes, the grid no longer exercises any control on the valve current, as the grid attracts positive ions during ionization. The valve current is limited only by the external circuit. When the anode voltage has dropped below 15 V, the valve deionizes and the grid regains control.

14. The mode of operation is then as follows :—The valve is initially biased negatively by the voltage drop

across R3 : condensers C2 and C3 charge up through R4 until they reach a potential sufficient to ionize the valve. When this occurs C2 and C3 discharge rapidly through V1, C1, R3B and L1, the latter serving to delay slightly the discharge time and hence to make the pulse developed across R3B roughly square-topped. This pulse is then fed to the grid of V2 via C4 and R1B and also down the three-way cable to synchronize the receiver time bases. With the switch S1 closed, the anode of V1 is grounded through R3B, and V1 becomes inoperative, the locking pulse then being provided from the A.A. No. 1 Mk. II transmitter through the three-way cable.



Circuit Reference	Description	Circuit Reference	Description
C 1	25 μ F 25 V working Electrolytic No. 4	R 9	100,000 Ω $\frac{1}{2}$ W No. 3A
C 2	.01 μ F 400 V working paper No. 3	R 10	2.2M Ω $\frac{1}{2}$ W
C 3	.005 μ F 400 V working paper No. 2	R 11	6,800 Ω $\frac{1}{2}$ W
C 4	.0002 μ F mica No. 1	R 12	22,000 Ω $\frac{1}{2}$ W
C 5	50 μ F ceramic No. 1	L 1	200 μ H choke
C 6	100 μ F ceramic No. 6	L 2	5.6 μ H choke
C 7	20 μ F ceramic No. 5	L 3	8 Turns osc. coil
C 8	3-36 μ F variable	L 4	20 H choke
C 9	.1 μ F 400 V working paper No. 5	L 5	Coupling coop
C 10	8 μ F 500 V working electrolytic No. 13	V 1	V.G.T. 128
C 11	410 μ F variable No. 1	V 2	V.R. 65
R 1	50 Ω $\frac{1}{2}$ W	V 3	V.R. 66
R 2	100,000 Ω 1 W No. 2A	V 4	5Z4G
R 3	1,000 Ω $\frac{1}{2}$ W No. 3A	N	Niphan plug N660 and Socket N662
R 4	220,000 Ω 1 W	G	Friction Drive 6.1 Reduction
R 5	10,000 Ω 1 W No. 2A	W	Westector W.6
R 6	220 Ω $\frac{1}{2}$ W No. 2A	T	Transformer 45V 50 c/s
R 7	1,500 Ω $\frac{1}{2}$ W	S. 1	Switch on-off No. 8
R 8	500,000 Ω Variable No. 7	M	Induction Motor

END

(51549/14/MG/ME-12C)

CRYSTAL RECTIFIERS

Characteristics, measured on different meters

[Based on W. O. (D. M. E.) T. I. No. T/M. 17]

EQUIPMENT

1. Crystal rectifiers.

DESCRIPTION

2. During recent measurements on the accuracy of Instruments, Testing, Radio Mechanic, Universal No. 1 (Triplett Meter 666-H), measurements of backwards and forwards resistance of nine different crystals were taken, on the 250,000 ohms range, and were compared with the resistances of the same crystals as measured with an Avometer 46 range, on the 100,000 ohms range. The results were as shown in Table 1.

3. It will be seen that both backwards and forwards, the resistance of the crystal measured on the Triplett meter is greater than that measured on the Avometer.

4. Fig. 1 shows simplified circuit diagrams of the two meters on the ranges used and it will be seen that both meters use the same voltage cell ($1\frac{1}{2}$ V) and that the Triplett meter has the greater series resistance. It follows that the voltage applied across the crystal will be greater in the case of the Avometer.

5. Fig. 2 shows a crystal characteristic curve, and the resistance measured is the reciprocal of the slope of the line joining the "working point" to the origin (0, 0.) The "working points" for the two meters are ($V_A, i_A 1$) and ($-V_A, -i_A 1$) for the Avometer and ($V_T, i_T 1$) and ($-V_T, -i_T 2$) for the Triplett meter where $V_A > V_T$. It will be clearly seen that the resistance measured by the Avometer in both directions is less than that measured by the Triplett meter. In the same way, changing the ohms range of either meter will change the "working point" and will again change the resistance measured.

ACTION

6. By all personnel concerned.

INSTRUCTIONS

7. When measuring crystal resistance with a meter, other than quoted in the publication giving the specification of the crystal, allowances must be made for the fact that different values of the resistance are to be expected.

Crystals	Avometer 46 Range on 100,000 ohms range			Triplett Meter 666-H on 250,000 ohms range		
	Backwards	Forwards	Ratio	Backwards	Forwards	Ratio
1	3,500 ohms	190 ohms	19.4	5,300 ohms	650 ohms	8.2
2	2,400 "	185 "	13.0	3,700 "	510 "	7.3
3	750 "	125 "	6.0	1,400 "	370 "	3.8
4	2,250 "	140 "	16.1	2,300 "	420 "	5.5
5	1,950 "	210 "	9.3	2,650 "	520 "	5.1
6	1,400 "	180 "	7.8	2,300 "	550 "	4.2
7	880 "	140 "	6.3	1,400 "	400 "	3.5
8	450 "	110 "	4.1	800 "	290 "	2.8
9	3,400 "	200 "	17.0	5,200 "	520 "	10.0

Table 1

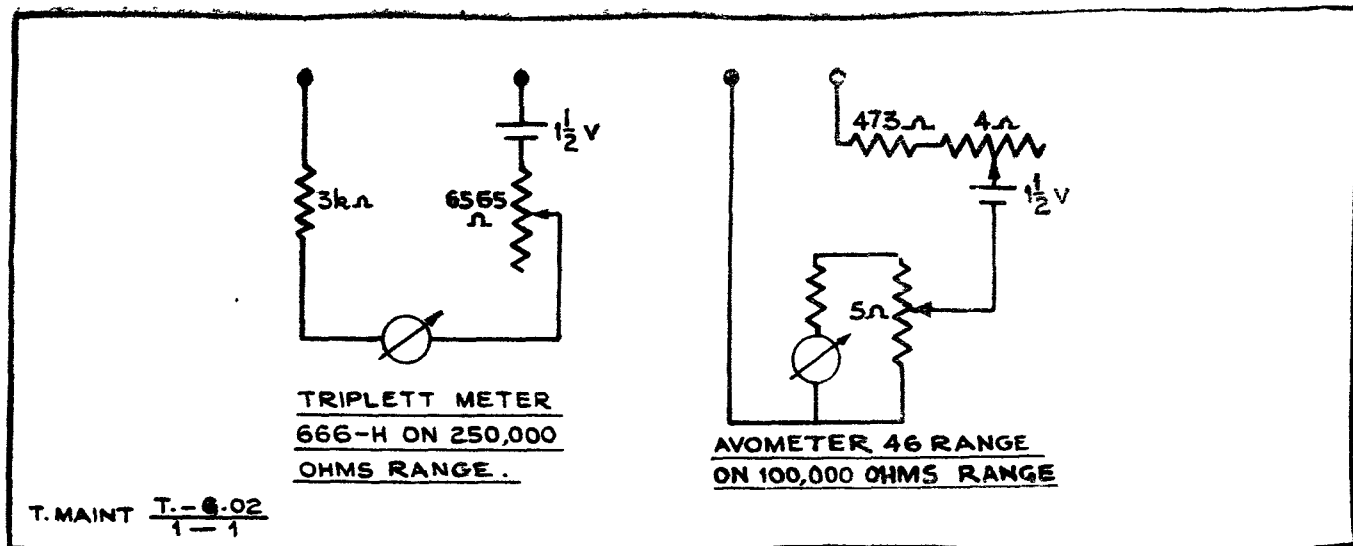


Fig. 1. Simplified circuit diagrams of the ranges used

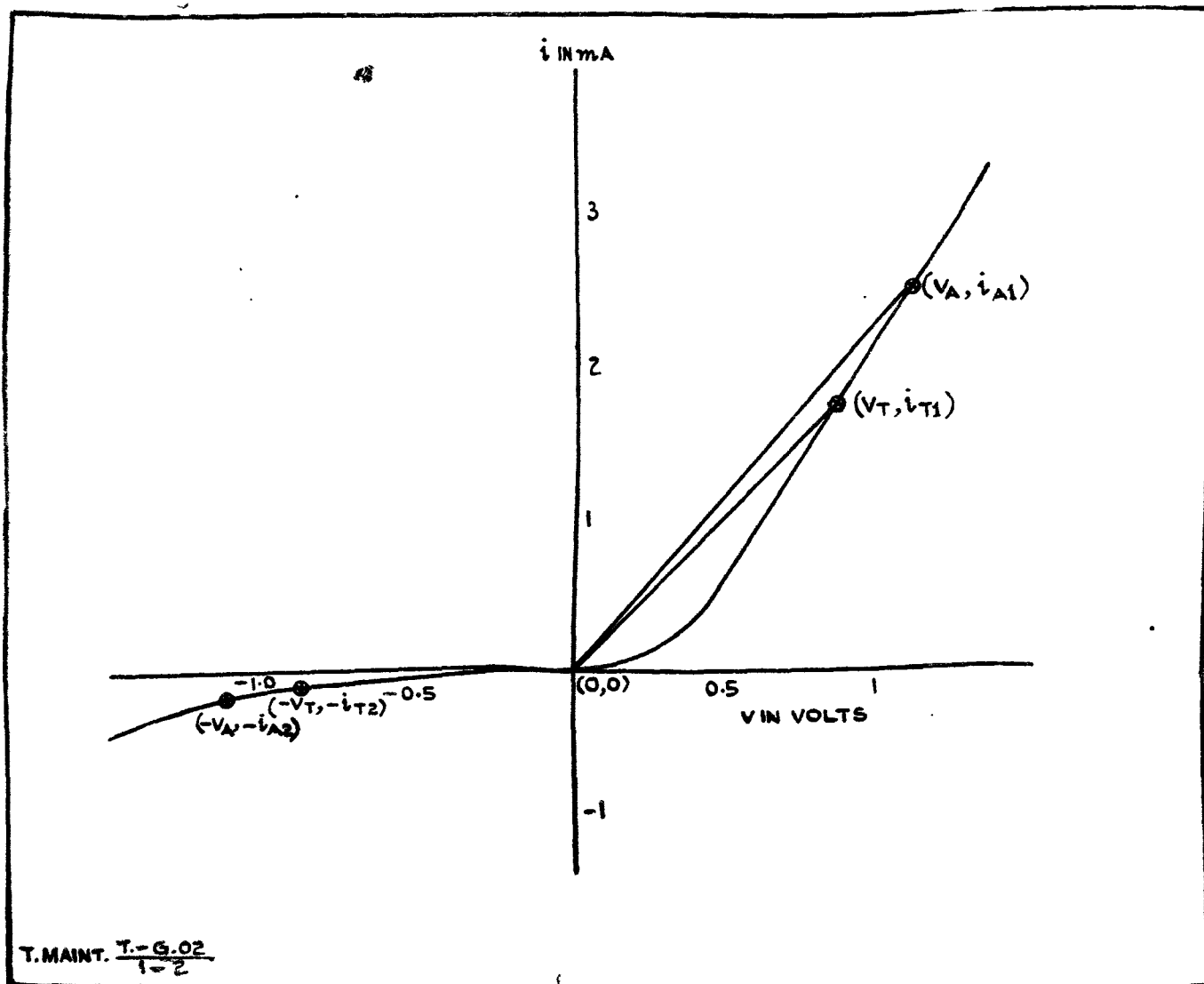


Fig. 2. Crystal characteristic curve
END

PRINCIPLES OF ELECTRICAL COMPUTING DEVICES.

Application of the varistor modulator circuit to servo systems.

(Based on E. M. E. R. Insts. and S/L E010/2)

1. The purpose of this circuit is to enable a constant source of A. C. voltage to be varied in amplitude by means of a varying D. C. modulating voltage and to be reversed in phase when the polarity of this D. C. voltage is reversed. This is used in

3. The copper oxide discs are similar to those commonly used for rectification, but to understand their use in this circuit they cannot be regarded simply as one-way current p. ths. A more detailed characteristic curve of their performance is shown by

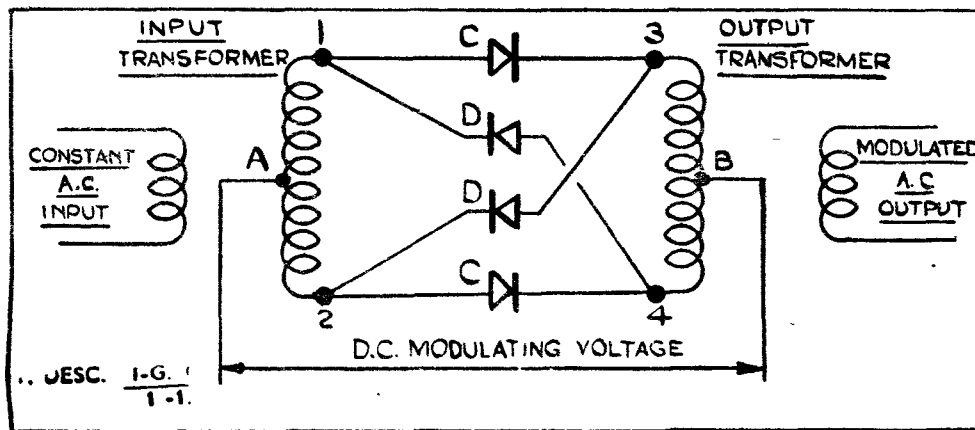


Fig. 1—Varistor modulator circuit.

servo systems to control one winding of a two-phase induction motor, the phase and amplitude of the voltage applied to the other winding being fixed.

2. The circuit is shown schematically in fig. 1. The constant A. C. source is connected to the primary of the input transformer, of which the secondary winding is centre-tapped. This secondary is, in turn, fed to the centre-tapped primary of the output transformer. Between the input and output transformers is connected a lattice network of copper oxide discs; the relative resistances of the four arms of this network will determine the magnitude and the phase of the voltage reaching the primary of the output transformer. These relative resistances are determined by the magnitude and polarity of the D. C. modulating voltage which is applied between the centre taps of the input and output transformers, the operation being as follows.

the full line curve in fig. 2 (in which, it should be noted, the scale of ohms is plotted logarithmically). It will be seen that when a D. C. voltage is applied across the disc in one direction, the disc presents a very high resistance to the current flow (over 100,000 ohms). As the D. C. voltage is reduced in magnitude to zero and then increased in reverse polarity, the resistance falls rapidly but not *instantaneously* to the very low value of only 10 or 20 ohms. This change of resistance will apply (approximately) not only to the D. C. voltage as shown but also to any small A. C. ripple that may be superimposed upon it, provided the A. C. amplitude is small compared with the total D. C. voltage sweep required to bring about the change of resistance (X-Y in fig. 2).

4. Referring again to the circuit diagram (fig. 1) if the point A is made positive relative to the point B, discs CC will have a low resistance and discs DD

DESCRIPTION
INST. G. 01

a high resistance and the output transformer will operate as though terminal 1 were connected to 3 and terminal 2 to 4. If, now, the D.C. voltage be reversed, discs DD become of low resistance and discs CC of high resistance. The result is as though

D. M. E. (INDIA)
TECHNICAL INSTRUCTIONS

increases in opposite polarity can clearly be followed by observing the combined effect of the plain and dotted curves (for the oppositely connected pairs of discs) of fig. 2. As the D.C. voltage is reduced, the A. C. output is also reduced until the two curves

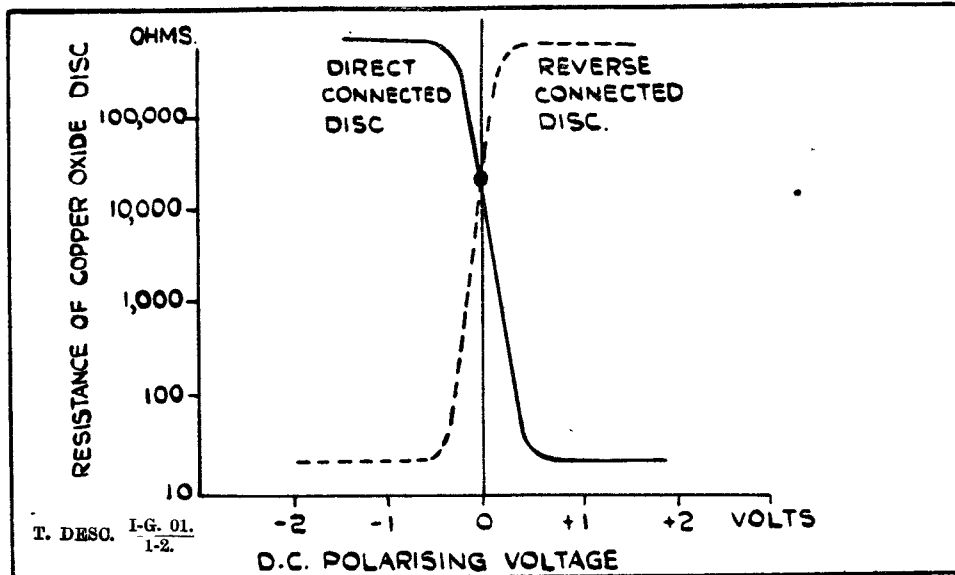


Fig. 2—Characteristic of copper oxide disc.

the connections between the input and output transformers had been interchanged and the phase of the output consequently reversed.

5. The complete cycle of operation as the D. C. voltage falls from a given value to zero and then

intersects at zero D.C. volts. At this point the resistance of the direct and reverse connected paths are equal and there is no A. C. output. As the D.C. increases again in opposite sign, the A. C. output will be built up in reversed phase.

END

(51241/5/ MG-ME 12)